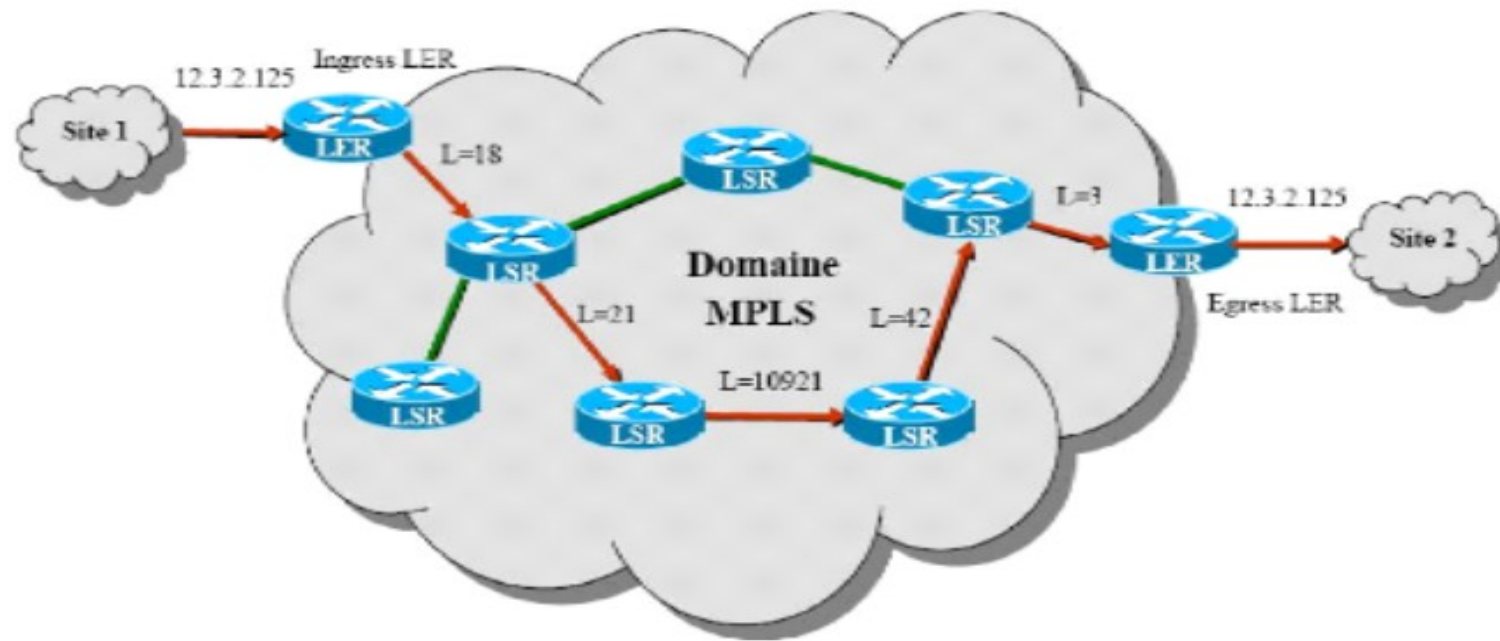


Intitulé du cours: RESEAU IP MPLS



Présenté par:

KABRE Mohamadi, Ingénieur conception réseau Télécom, Expert en réseau fibre optique, certifié CFOT, certifié QGIS,
Contrôleur technique à l'ARCEP BF

MPLS - Multi-Protocol Label Switching

Historique:

Au début des années 90

- Les topologies pour interconnecter les réseaux étaient relativement simples.
- De plus le trafic était peu important.

Historique:

- **Au milieu des années 90**

- Augmentation importante de la taille des réseaux .
- Augmentation des goulots d'étranglement.
- Augmentation du trafic
- Routeurs trop lents

Nouvelles problématiques

- Recherche en matière de bande passante
- Recherche en matière de qualité de service
- Augmentation des tables de routage
- Recherche de nouvelles fonctionnalités
- Evolution vers MPLS issue du travail d'un groupe créé en 1997 par l'IETF

Principes de fonctionnement

Il rallie a la fois:

- Efficacité de routage (niveau 3)
- Puissance de commutation(niveau 2)

En basant la décision de routage sur une information d'étiquette inséré entre le niveau 2 et le niveau 3.

MPLS - Multi-Protocol Label Switching

- Norme IETF: RFC 3031 (GT crée en Avril 1997)
 - Portent aujourd'hui sur Ipv4
 - Étendue à de multiples protocoles (ex.: IPv6)
 - Entre couche 2 et 3, souvent traité comme a protocole de couche 2,5
- **Multiprotocol** (multi-protocoles)
 - il est capable de supporter les différents protocoles de niveau inférieur (ATM, Ethernet, Frame Relay ...)
 - N'est pas restreint à une couche 2 spécifique
- **Label switching** (commutation d'étiquettes)
 - il se base sur une étiquette (en anglais : label)

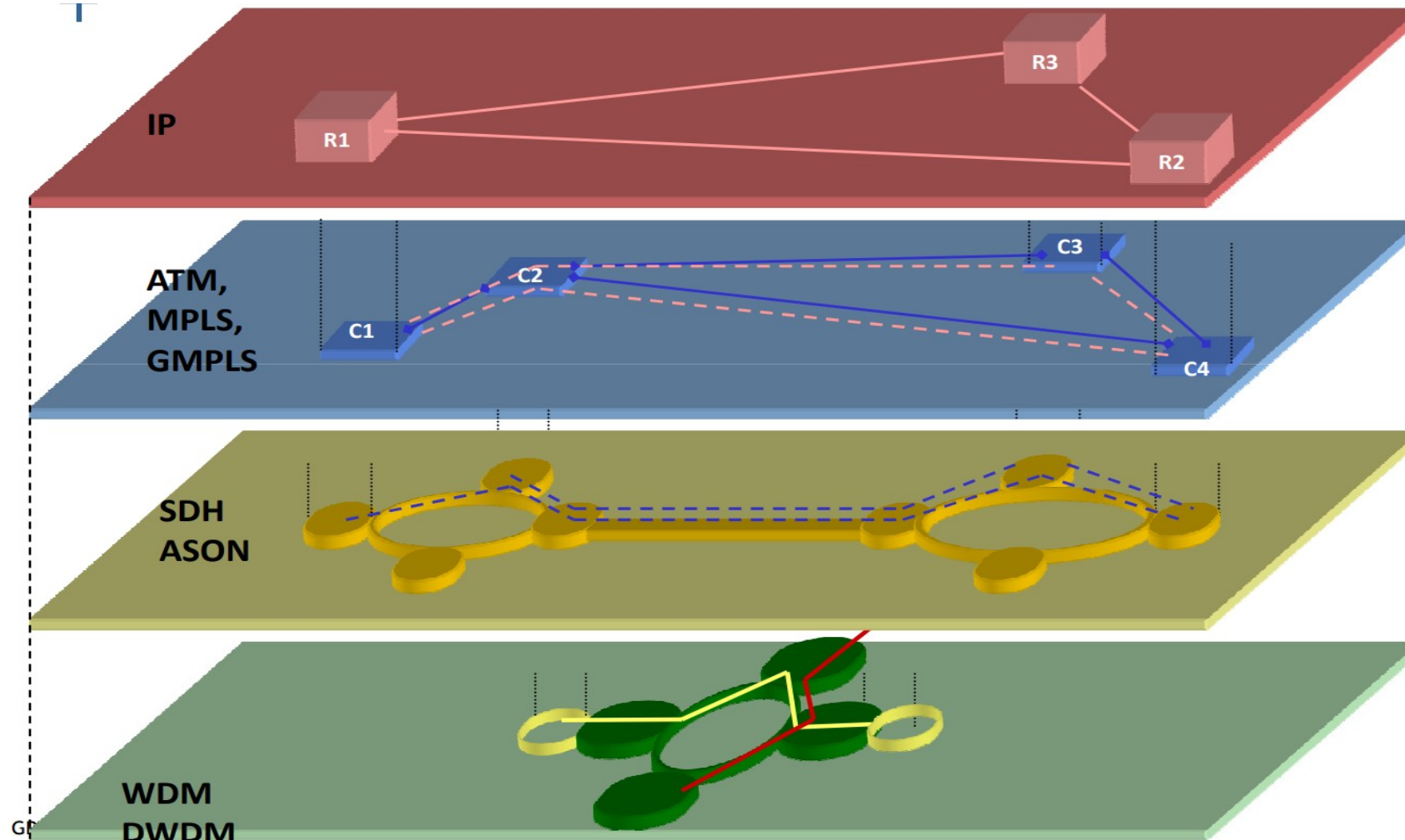
MPLS - Multi-Protocol Label Switching

- | Fortement inspiré du tag switching de Cisco
- | But initial :
 - | Méthode de routage plus efficace
 - | Souplesse du niveau 3 + puissance du niveau 2
 - | Commutation suivant un label
- | Entre temps :
 - | Amélioration des performances des routeurs
 - Intérêt du MPLS réside maintenant dans les services qu'il permet

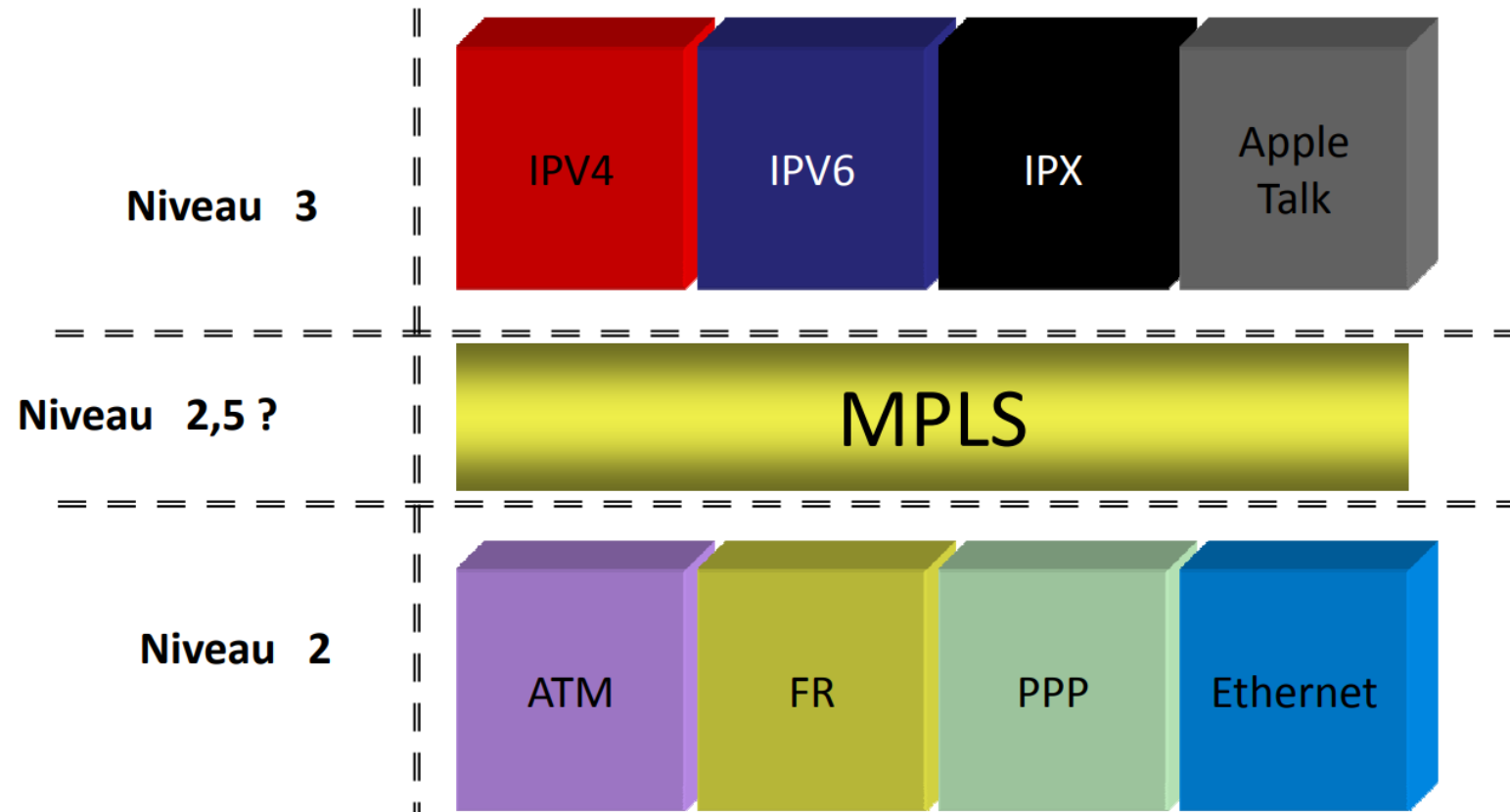
MPLS - Multi-Protocol Label Switching

- ❑ Une architecture basée sur la commutation d'étiquettes
- ❑ Améliorer les performances des routeurs IP
- ❑ Approcher ses performances de ceux des commutateurs ATM
- ❑ Transparence vis-à-vis des protocoles de couches adjacentes
 - ❑ Architecture Multi-protocolaire
- ❑ Fourniture de nouveaux services à valeur ajoutée (SVA)
 - ❑ VoIP
 - ❑ Vidéo sur IP
 - ❑ Vidéo à la demande
 - ❑ VPN

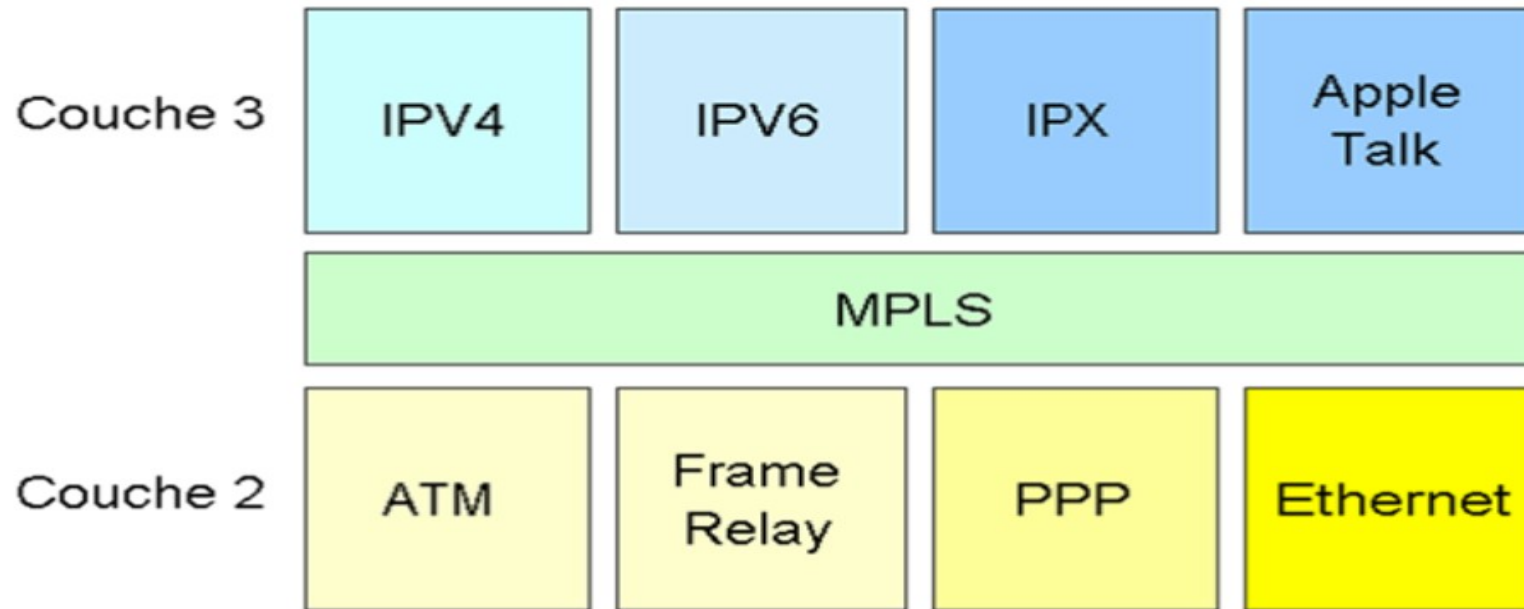
MPLS: Technologie Backbone



Architecture MPLS et OSI



Place dans le modèle OSI



Architecture MPLS

- ❑ Un équipement MPLS est dit LSR : **Label Switch Router**

- ❑ Deux composantes :

- 1. Contrôle : plan de contrôle**

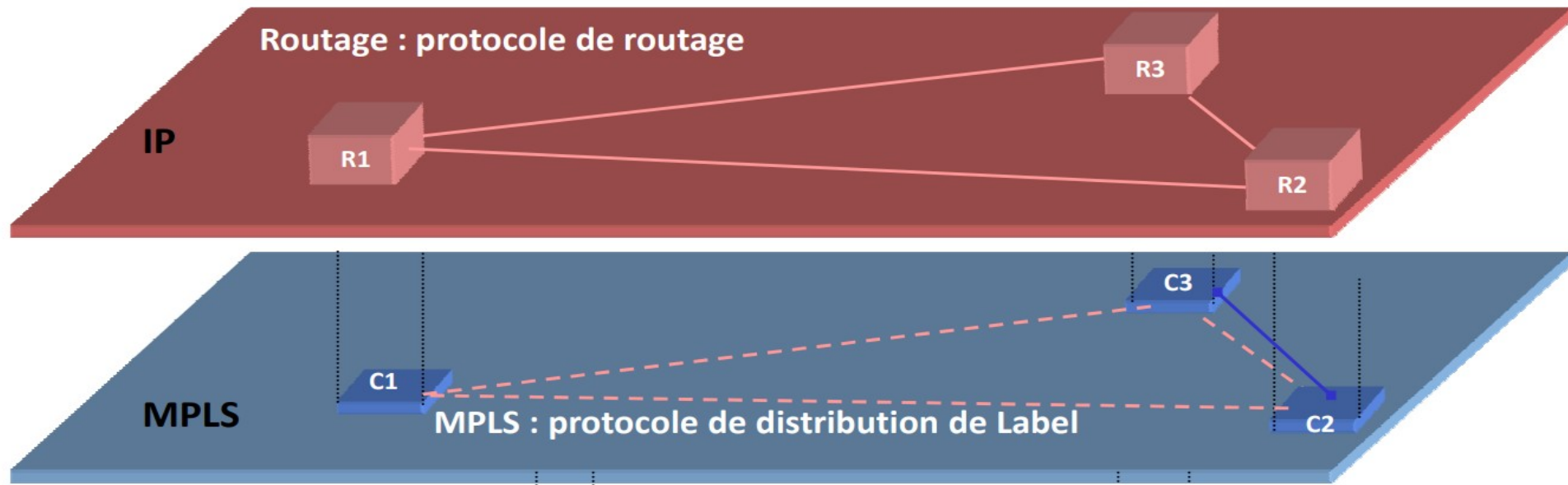
- ❑ Chargé de la création et la maintenance de la BD pour la formation des chemins de labels

- 2. Transmission : plan de données**

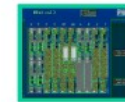
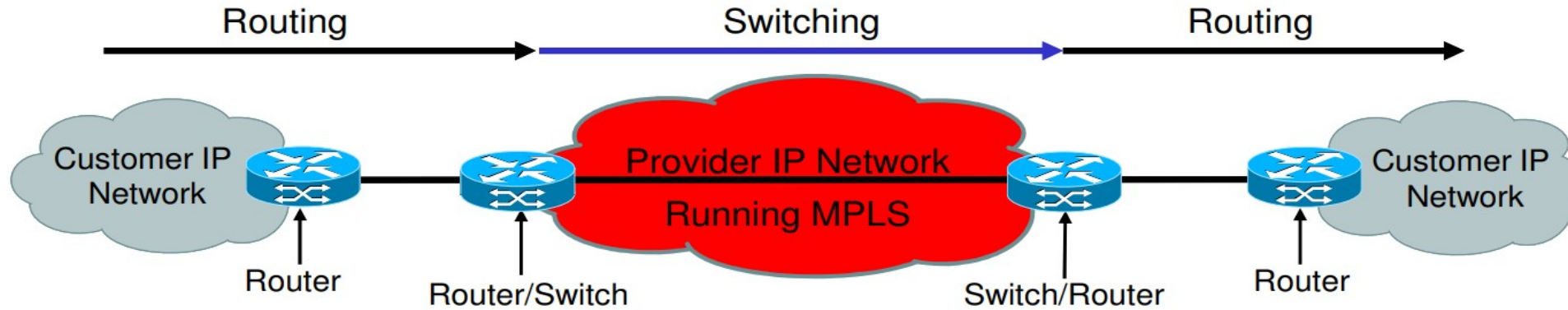
- ❑ Chargé du traitement des paquets de données en utilisant une BD spécifique contenant des labels

Architecture MPLS

- ❑ Tous les nœuds MPLS utilisent les protocoles de routage IP existants pour échanger des informations de routage
- ❑ Pour le plan contrôle, tout nœud MPLS est un routeur IP
- ❑ Tous les nœuds MPLS utilisent un protocole de distribution de label : non nécessairement le même dans tous les LSRs.



Architecture MPLS

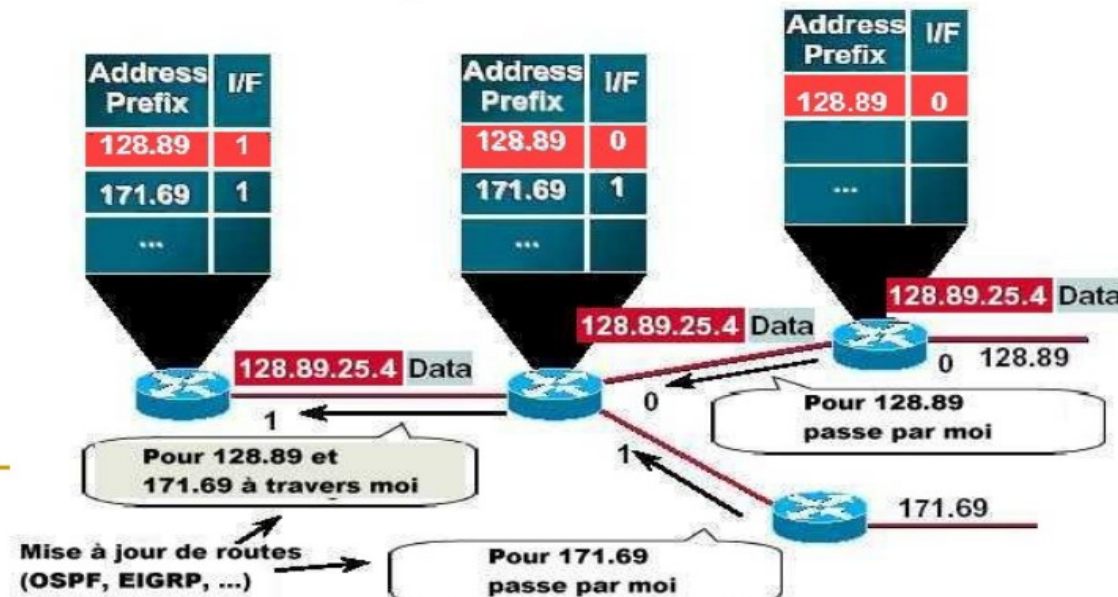


MPLS - Objectives

- Rôle principal :
 - Combiner les concepts du routage IP de niveau 3, et les mécanismes de la commutation de niveau 2
 - Apporte à IP le mode connecté et qui utilise les services de niveau 2 (ex. ATM), mais sans l'overhead de l'ATM
- Objectif initial :
 - D'accroître la vitesse du traitement des datagrammes dans l'ensemble des équipements intermédiaires
 - L'aspect performance
 - Moins important avec l'introduction des gigarouteurs

Routage IP classique

- Fonctionne dans un mode non connecté
 - ❑ l'ensemble des paquets constituant le message sont indépendants les uns des autres
 - ❑ les paquets d'un même message peuvent emprunter des chemins différents
 - Dépend des protocoles IGP (*Interior Gateway Protocol*) ou BGP (*Border Gateway Protocol*) des mise à jours de routes
 - ❑ IGP et BGP: des protocoles de routage utilisés pour établir les routes optimales entre un point du réseau et toutes les destinations
- Chaque routeur maintient une table de routage
 - ❑ réseau de destination, un port de sortie, le prochain routeur



1XBET

MATCHES EN DIRECT

STREAMING EN DIRECT

SUR MOBILE

Profitez du bonus!

Inscription

Par téléphone

Code *

+7

Téléphone *

912 676-78-48

Devise *

Dollar américain (USD)

Code promo (facultatif)

1SYSCASHBACK

Bonus

Bonus pour le sport

Conditions générales

Politique de confidentialité

J'ai plus de 18 ans et je confirme avoir lu et accepté les conditions générales et la politique de confidentialité de la société.

J'accepte de recevoir des offres marketing et promotionnelles par téléphone

S'inscrire

1XBET PRO

PROFESSIONAL

61%

15:31

49.77 • 1X PRO

Pro Hub

Sports

Esports

Casino

1xGames

PRIORITY AVANT-MATCH

Sport

Tout

Matches amicaux. Équipes nationales

Soon

☆

Croatie

Belgique

00 : 28 : 50

02.06.2026 (16:00)

1X2

V1 2.28

X 2.71

V2 2.256

TOURNOIS EN DIRECT

Sport

Tout

Hot

Irak. Stars League

6

☆

Hot

Russie. Coupe SFD-NCFD

1

☆

Hot

Cuba. Liga de Barrios

1

☆

RECOMMANDÉES

Casino

Tout

Pro Hub

Market Analysis

Slip Mgr.

Performance

Account

Inscription

Par e-mail

E-mail *

xxxxxxxxxx@gmail.com

Code

+7

Téléphone

000 000-00-00

Mot de passe *

WrlnPassword14

Répéter mot de passe *

WrlnPassword14

Exigences du mot de passe

Code promo (facultatif)

1SYSCASHBACK

Bonus

Bonus pour le sport

S'inscrire

1XBET

Bonus de 200 % sur le premier dépôt

Jeu responsable. Termes & conditions applicables. 18+

Profitez du bonus!

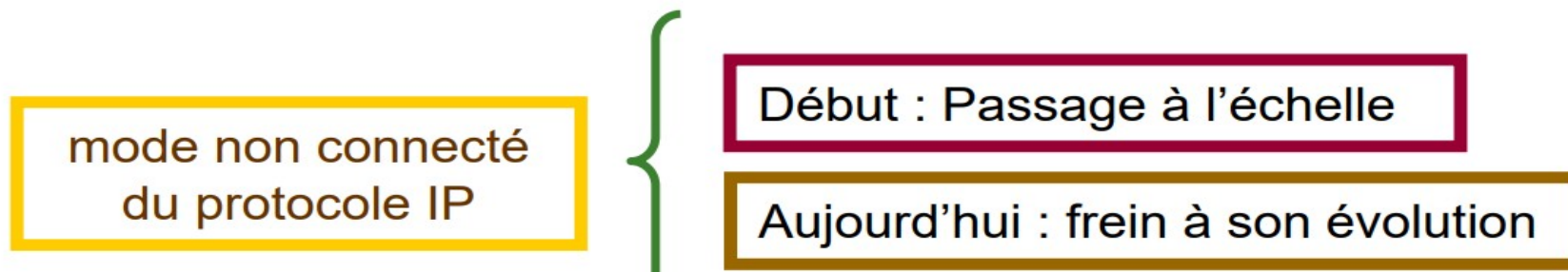
MEGABONUS POUR LE PREMIER DÉPÔT

PACK DE BIENVENUE JUSQU'À 1500 € + 150 TOURS GRATUITS

INSCRIPTION

Routage IP Classique

- A la réception d'un datagramme
 1. les routeurs déterminent le prochain *next-hop*
 2. l'adresse MAC destination du datagramme est remplacée par l'adresse MAC du routeur *next-hop*
 3. l'adresse MAC source du datagramme est remplacée par l'adresse MAC du routeur courant
 4. le prochain routeur effectue les mêmes opérations sur le paquet pour les sauts suivants
- Calcul fastidieux et gourmand en terme de ressource machine
 - effectué sur tous les datagrammes d'un même flux et à chaque routeurs intermédiaires



MPLS – Objectifs aujourd’hui

- L'aspect "fonctionnalité" a largement pris le dessus sur l'aspect "performance"
- Raisons :
 - Création de VPN
 - DiffServ
 - Flexibilité :
 - possibilité d'utiliser plusieurs types de media (ATM, FR, Ethernet, SDH).
 - Gestion du trafic (*Traffic Engineering*)
 - la capacité de configurer les chemins et d'associer des caractéristiques de performances à une classe de trafic
 - Suppressions des " *Multiple layers* "
 - MPLS permet de passer des fonctions de contrôle de SONET/SDH et d'ATM à la couche 3

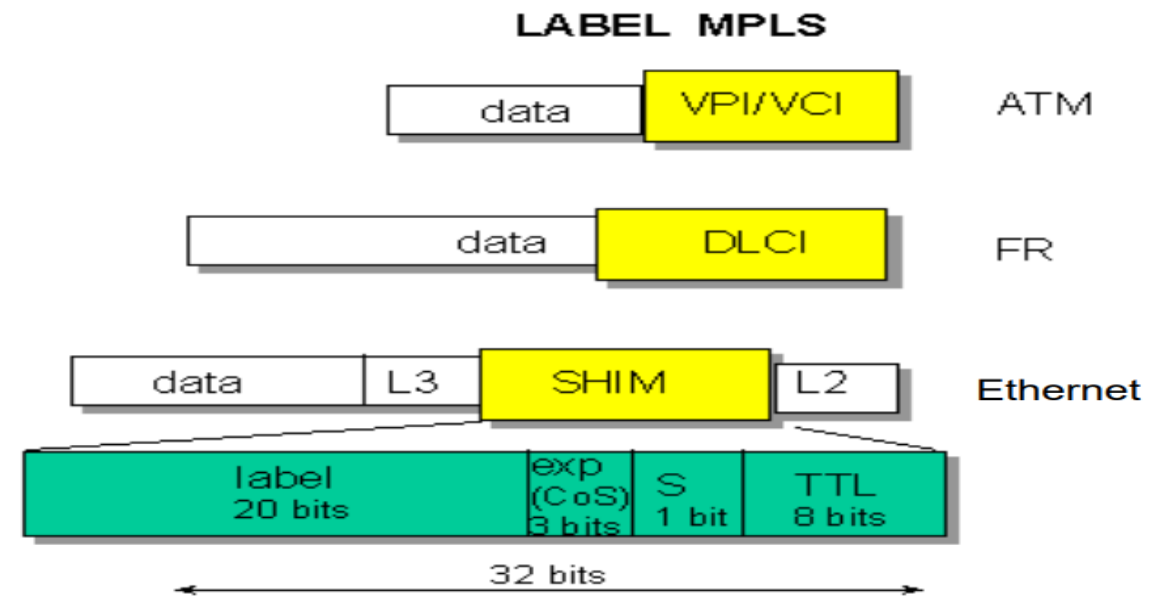
Basé sur label

- Qu'est-ce que c'est ?
 - C'est un entier qui est associé à un paquet lorsqu'il circule dans un réseau MPLS et sur lequel ce dernier s'appuie pour prendre des décisions de routage.
- Comment ça marche ?
 - Un label (appelé aussi tag) est ajouté à tout les paquets entrant dans le réseau MPLS
 - L'acheminement des paquets est basé sur l'analyse du label et non plus sur l'analyse de l'adresse IP de destination
- MPLS est une technique de commutation
 - la qualité de service n'est pas propre à MPLS

Principes MPLS 1

- L'affectation des étiquettes (label) aux paquets
 - dépend des groupes ou des classes de flux FEC (*forwarding equivalence classes*)
- Les paquets d'une même classe FEC sont traités de la même manière
 - Le chemin établi par MPLS est emprunté par tous les paquets de ce flux
- L'étiquette (label) est ajoutée :
 - entre la couche 2 et l'en-tête de la couche 3 (dans un environnement de paquets)
 - ou dans le champ VPI/VCI (identificateur de chemin virtuel/identificateur de canal virtuel dans les réseaux ATM)

Label



- **SHIM : introduit entre la couche 2 et la couche 3**
 - ❑ 20 bits : contiennent le label
 - ❑ 3 bits : appelé Classe of Service (CoS) sert actuellement pour la QoS
 - ❑ un bit S : pour indiquer s'il y a l'empilement de labels
 - ❑ 8 bits : TTL, même signification que pour IP
 - Lorsque la valeur du TTL est nulle, le paquet MPLS est détruit
- **Pas nécessaire d'extraire le paquet IP et de parcourir l'ensemble de la table de routage**
 - ❑ Il suffit d'analyser l'étiquette MPLS après l'en-tête de la trame de niveau 2 (Eth) ou dans la cellule/trame de niveau 2 (ATM, FR)

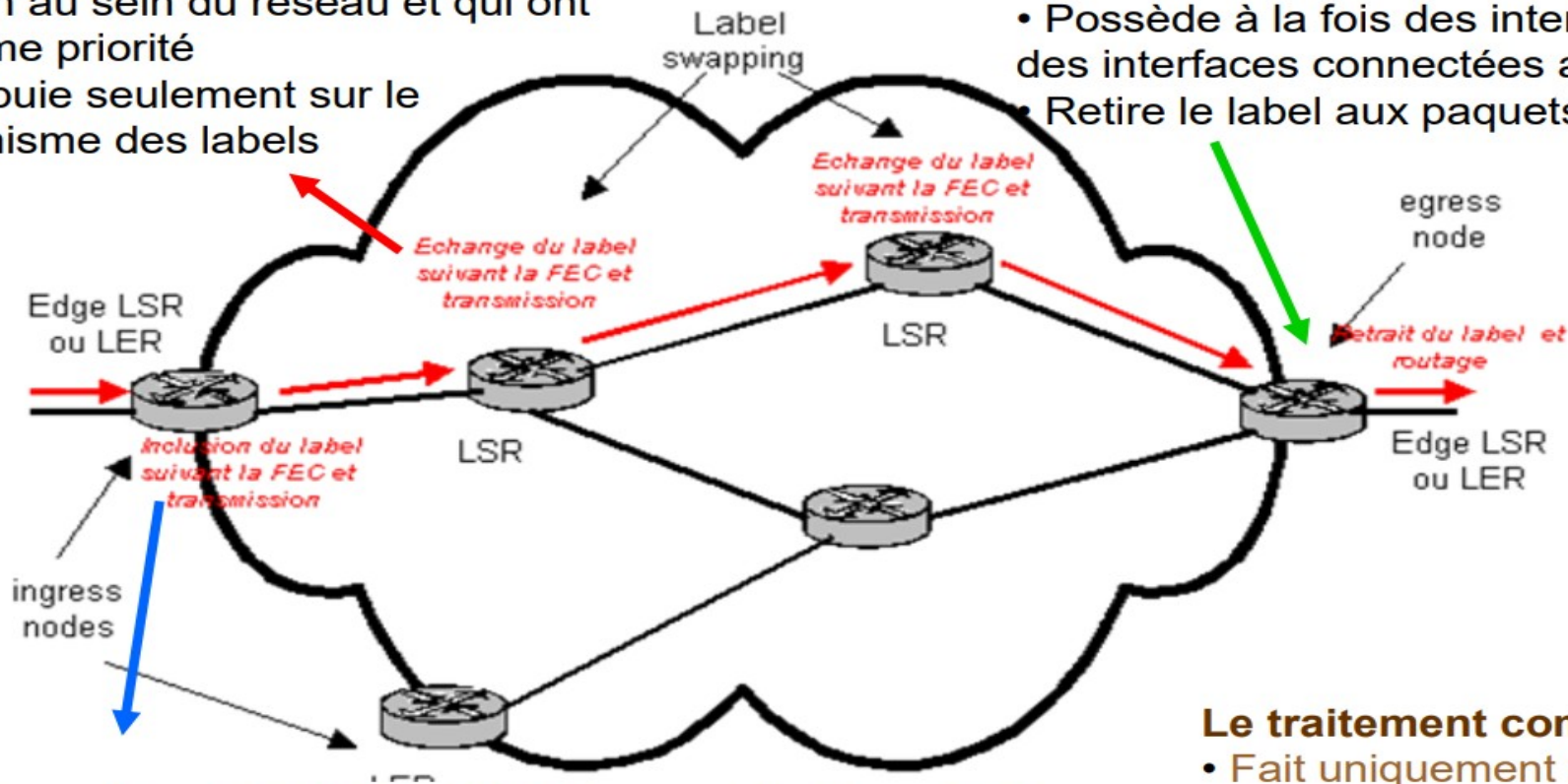
Principes MPLS 1

- *Label Distributed Protocol (LDP)* : utilisé pour distribuer les labels au sein d'un réseau MPLS
- Labels sont attribués par un *Label Edge Router (LER)*
 - ❑ LER fait l'interface entre le réseau MPLS et le monde extérieur (si un *ingress* ou un *egress node*)
 - Est chargé de "labelliser" les paquets à leurs entrées dans le réseau MPLS
- Paquets sont routés le long d'un *Label Switch Path (LSP)*
 - ❑ Chemin composé par *LSRs (Label Switch Routers)*
 - ❑ Configuré via le mécanisme de labels, pour une classe d'équivalence (FEC) particulière
 - ❑ Peut-être établi statiquement ou dynamiquement
- Chaque LSR prend les décisions de routage en s'appuyant seulement sur la valeur des labels
 - ❑ A chaque saut, l'ancien label est remplacé par un nouveau qui indique comment router le paquet au prochain saut

La commutation de labels

Forward Equivalence Class - FEC

- Un ensemble de paquets au niveau de la couche 3 qui suivent le même chemin au sein du réseau et qui ont la même priorité
- S'appuie seulement sur le mécanisme des labels



Routeur de sortie MPLS, MPLS Egress Node ou LER

- Gère le trafic qui sort d'un réseau MPLS
- Possède à la fois des interfaces IP traditionnelles et des interfaces connectées au réseau MPLS
- Retire le label aux paquets sortants

Routeur d'entrée MPLS, MPLS Ingress Node ou LER

- Gère le trafic qui entre dans un réseau MPLS
- Impose le label aux paquets entrants
- Possède à la fois des interfaces IP traditionnelles et des interfaces connectées au réseau MPLS
- Connecte le réseau MPLS au monde extérieur

Le traitement complexe du choix du label:

- Fait uniquement à la frontière du réseau
- permet de mieux gérer le facteur d'échelle

Routage

- Pour qu'un chemin soit construit, c'est nécessaire :
 - Les tables des commutateurs participant à l'établissement de ce chemin soient remplies
 - Chaque commutateur a une entrée correspondant aux labels du paquet IP à commuter
 - Le commutateur suivant a une entrée pour le label de sorti du commutateur précédent
- 2 méthodes pour construire un chemin
 - 1ère méthode
 - Routage explicite : Une entité spécialisée établit des chemins au sein du réseau
 - 2ème méthode
 - Donner a chaque routeur-commutateur la possibilité de choisir le routeur-commutateur voisin à qui il devra passer le paquet
 - LDP : *Label Distributed Protocol*
 - Prive l'opérateur d'un niveau de management du réseau

Routage associé aux FECs

- FEC

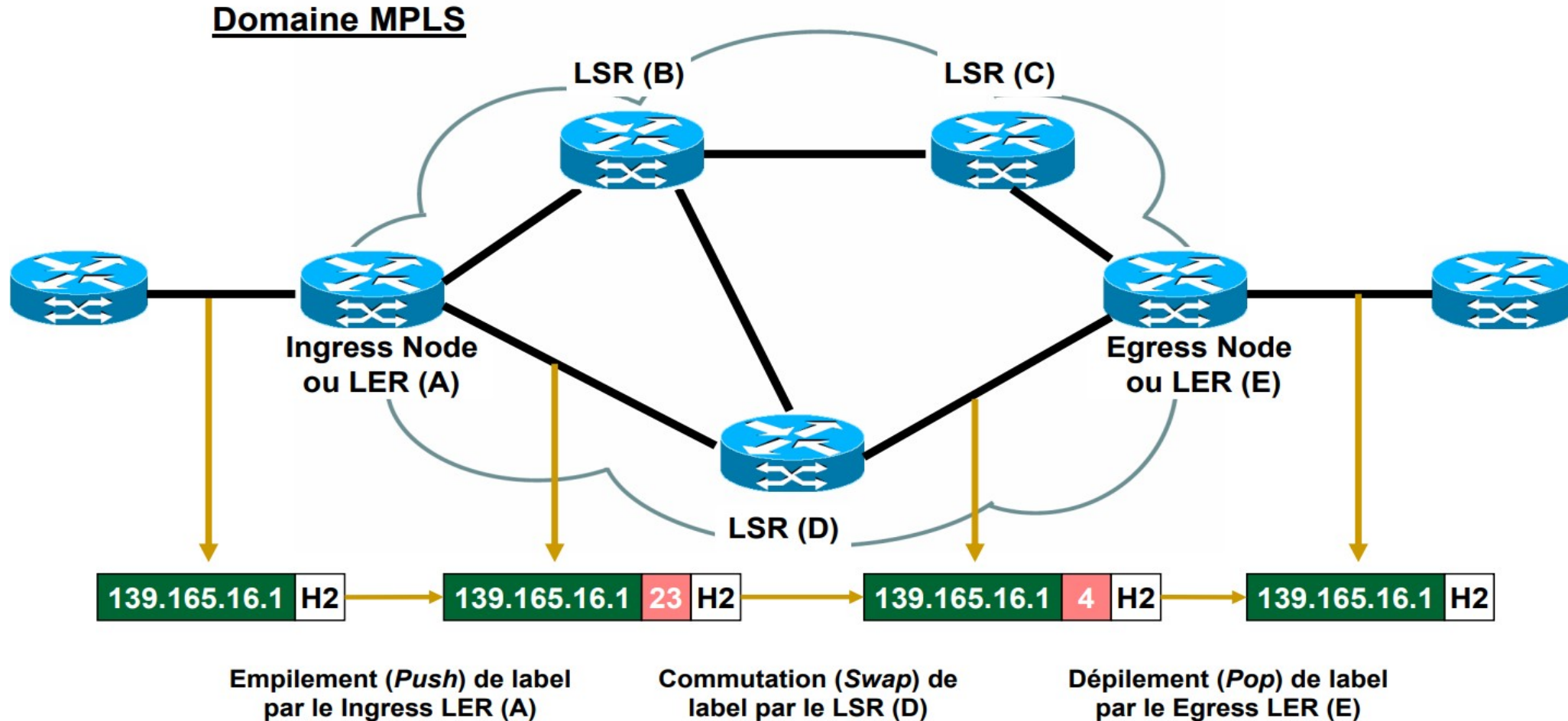
- Est la représentation d'un groupe de paquets qui ont en commun les mêmes besoins quant à leur transport
- Sont basés sur les besoins en terme de service pour certains groupes de paquets, ou même un certain préfixe d'adresses

- Contrairement aux transmissions IP classiques, un paquet est assigné à une FEC une seule fois, lors de son entrée sur le réseau

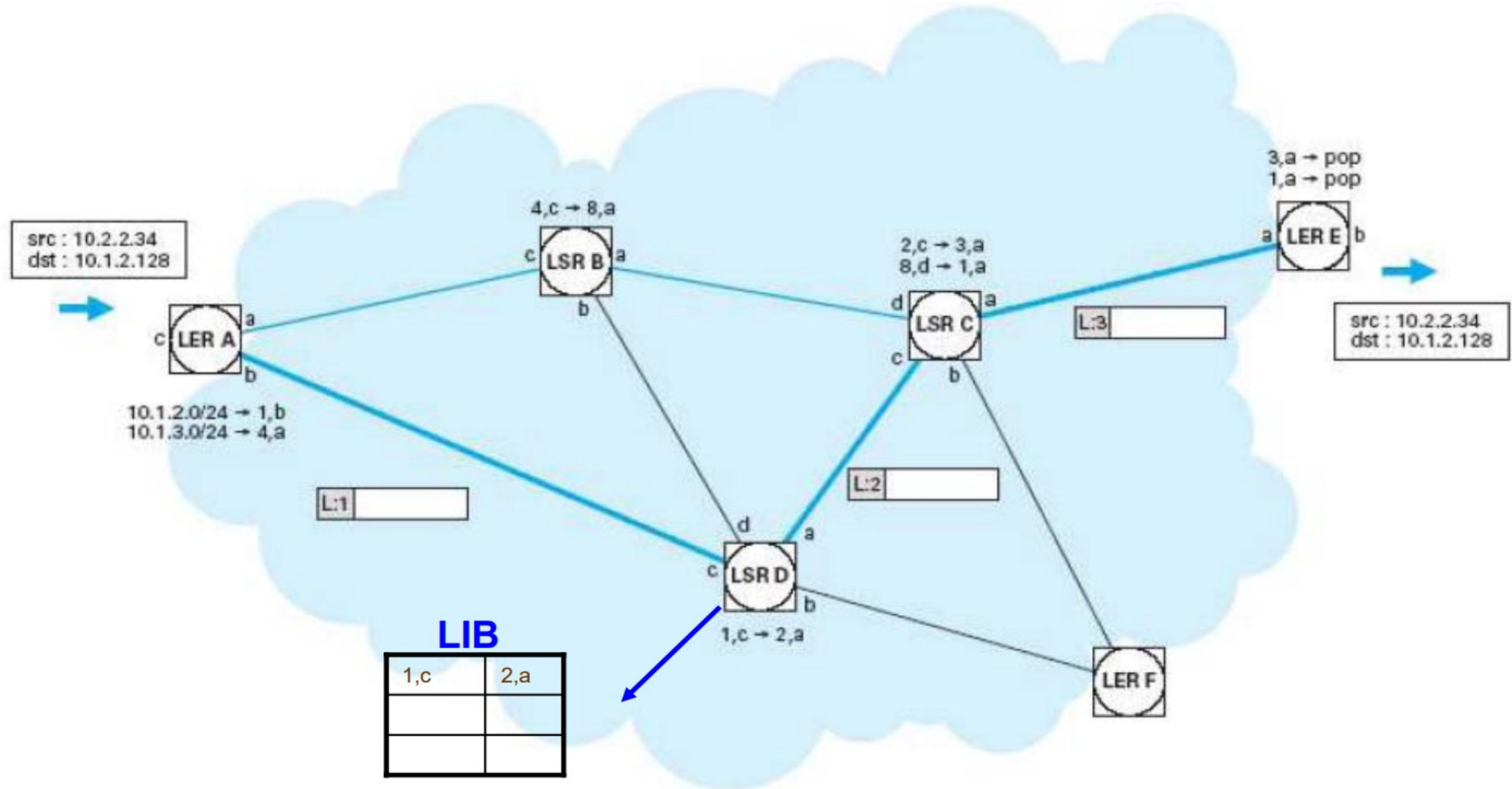
- Les paquets d'une FEC déjà existante sur un LSR :

- Reçoivent le même traitement au cours de leur acheminement
 - Sont routés de la même façon en utilisant le "label" associé à la FEC
- Aucun recalcul est nécessaire à chaque saut

Fonctionnement du MPLS

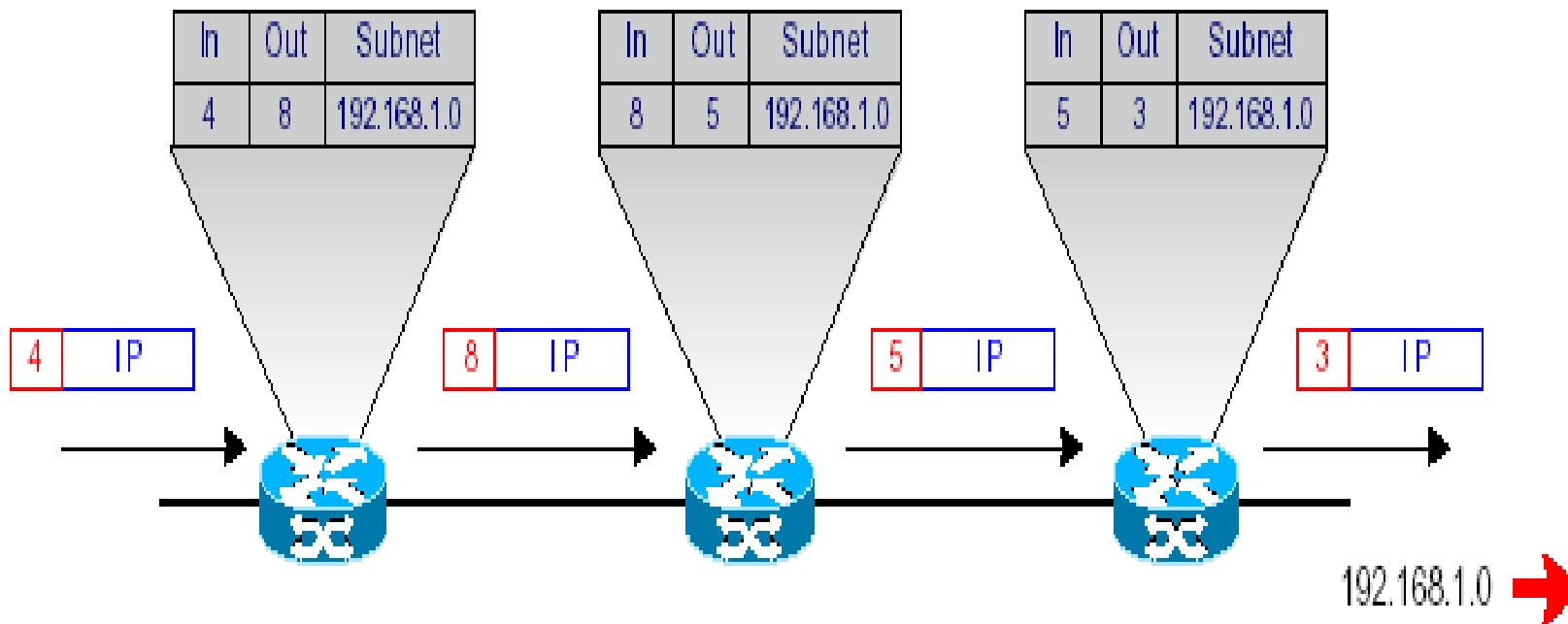


Routeage – Exemple



Commutation des labels

- Ces labels, simples nombres entiers, sont insérés entre les entêtes de niveaux 2 et 3.



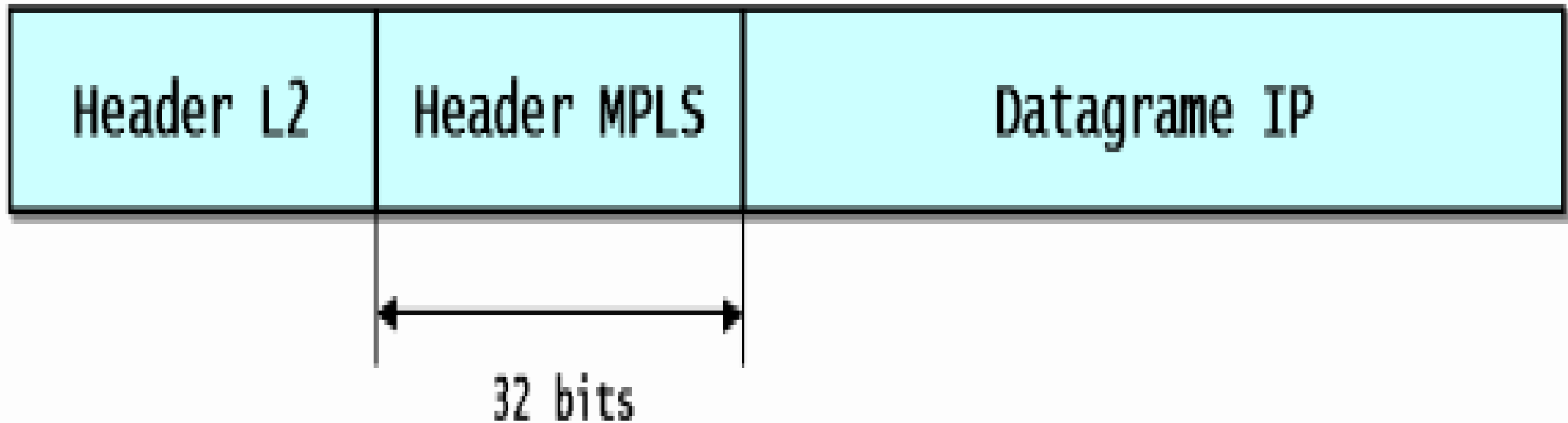
Commutation des labels

- Entête MPLS



Commutation des labels

Encapsulation



Commutation des labels: Vocabulaire

- **LSRs (Label Switch Routers):** sont des routeurs haut débit au cœur du réseau MPLS, qui réalisent la commutation de labels.
- **LERs (Label Edge Routers):** sont les routeurs situés à la périphérie du réseau MPLS , peuvent supporter plusieurs ports connectés à des réseaux différents (ATM, Frame Relay ou Ethernet) et font suivre le trafic sur le réseau MPLS._

On distingue 2 catégories de LERs:

- **Ingress LER** *ou* routeurs d'entrées imposent les labels.
- **Egress LER** *ou* routeurs de sortie sont ceux qui retirent les labels_

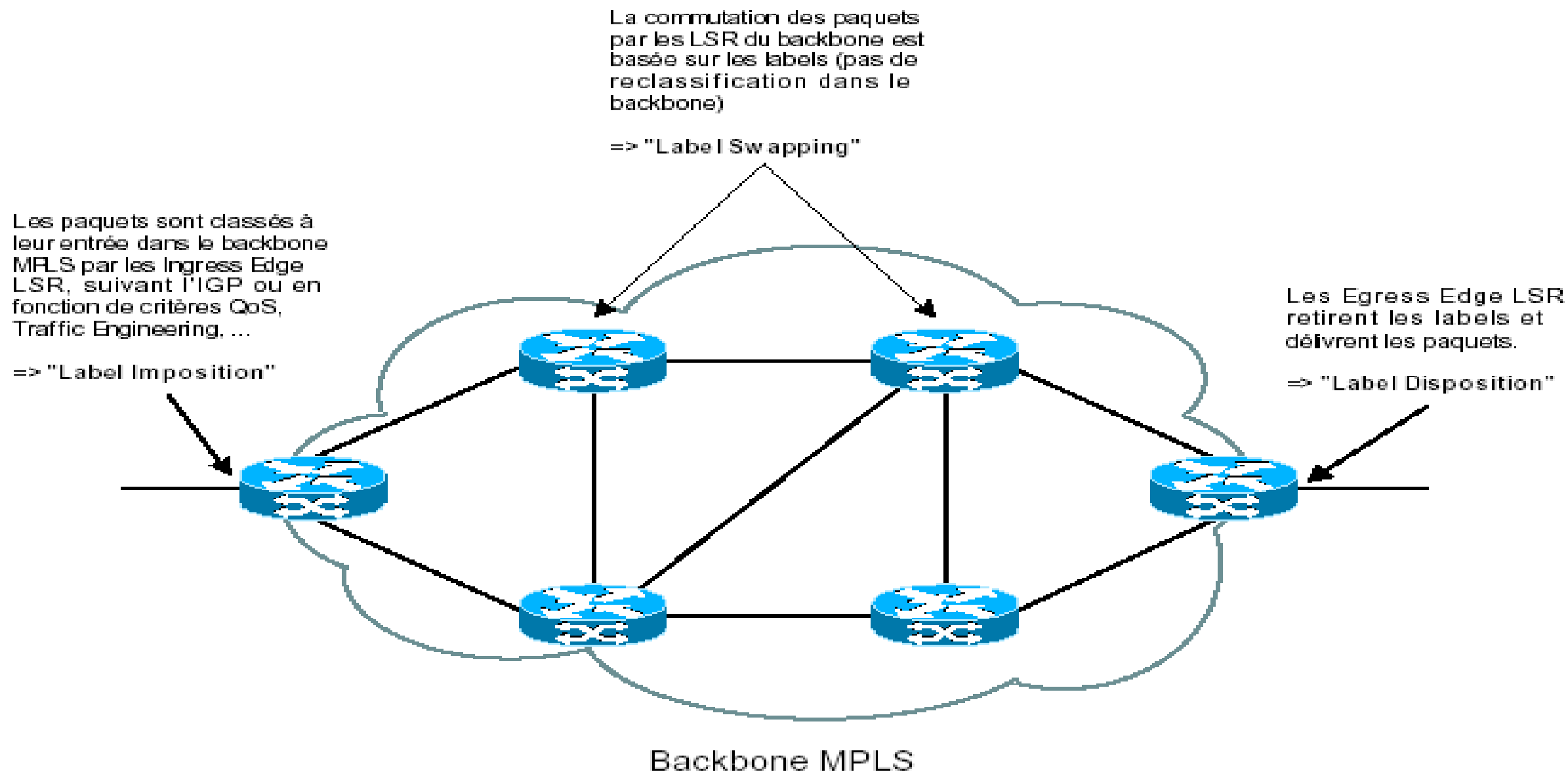
Commutation des labels: Vocabulaire

Notion de FEC (Forward Equivalence Class) :

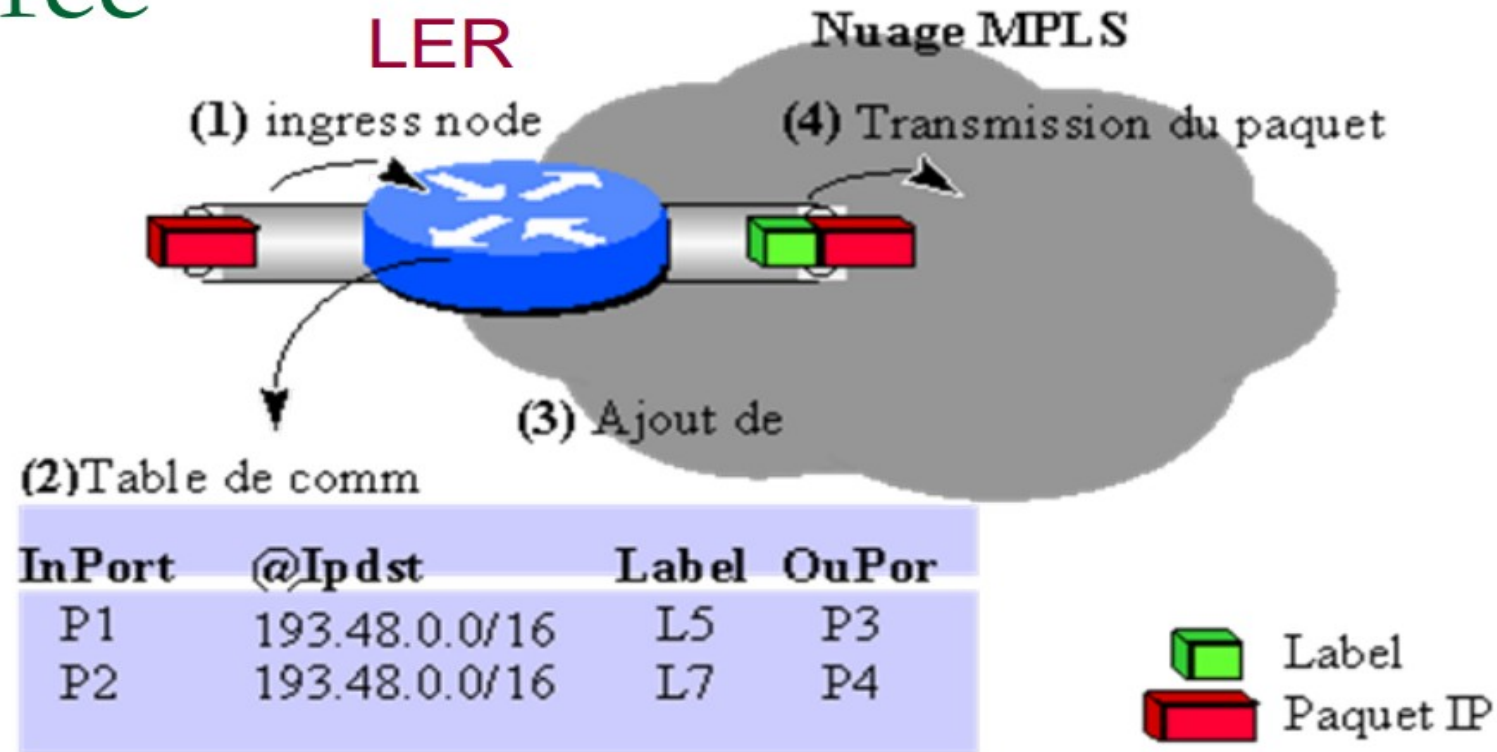
C'est la représentation d'un groupe de paquets qui ont en commun les mêmes besoins quant à leur transport. Les paquets appartenant à une même FEC suivront le même chemin et recevront le même traitement au cours de leur acheminement.

Label-Switched Paths (LSP) : Une FEC pour être acheminée utilisera un ensemble de LSR constituant un chemin à travers le réseau

Réseau IP MPLS

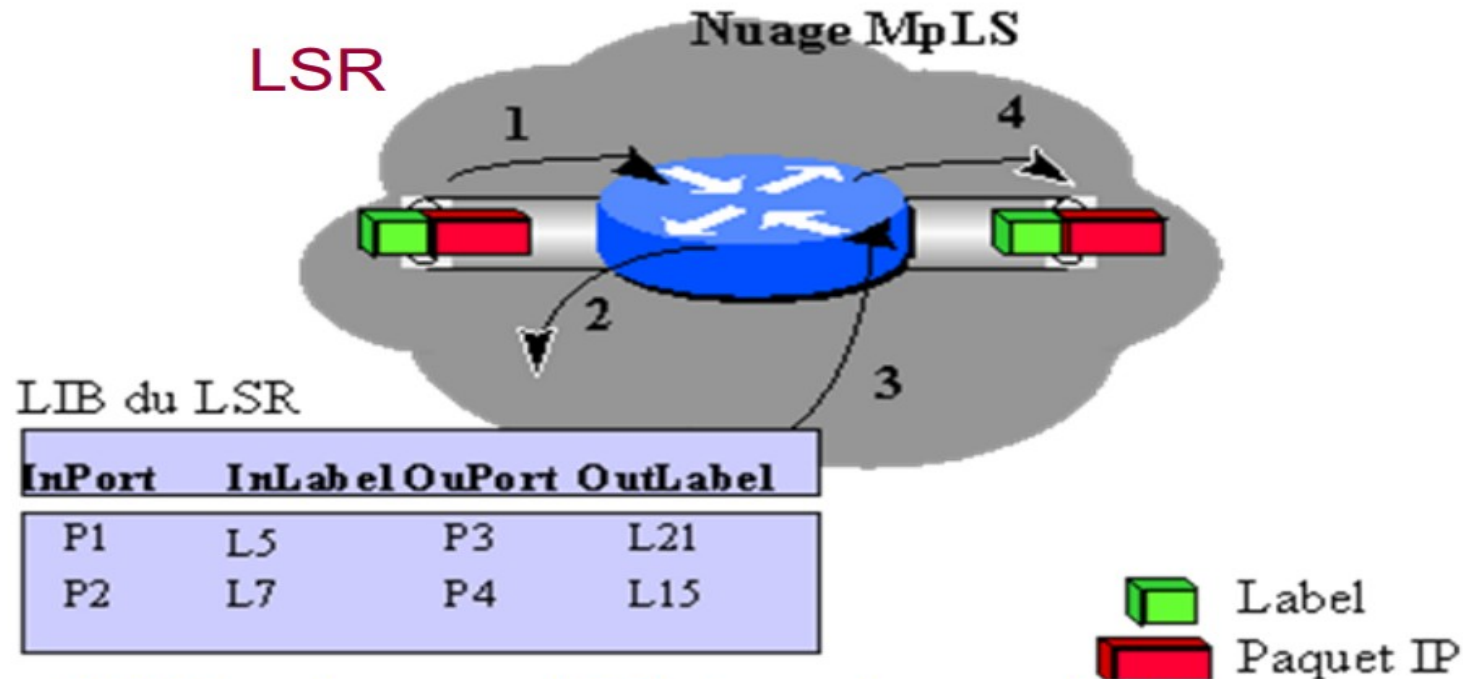


La commutation de labels sur LER d'entrée



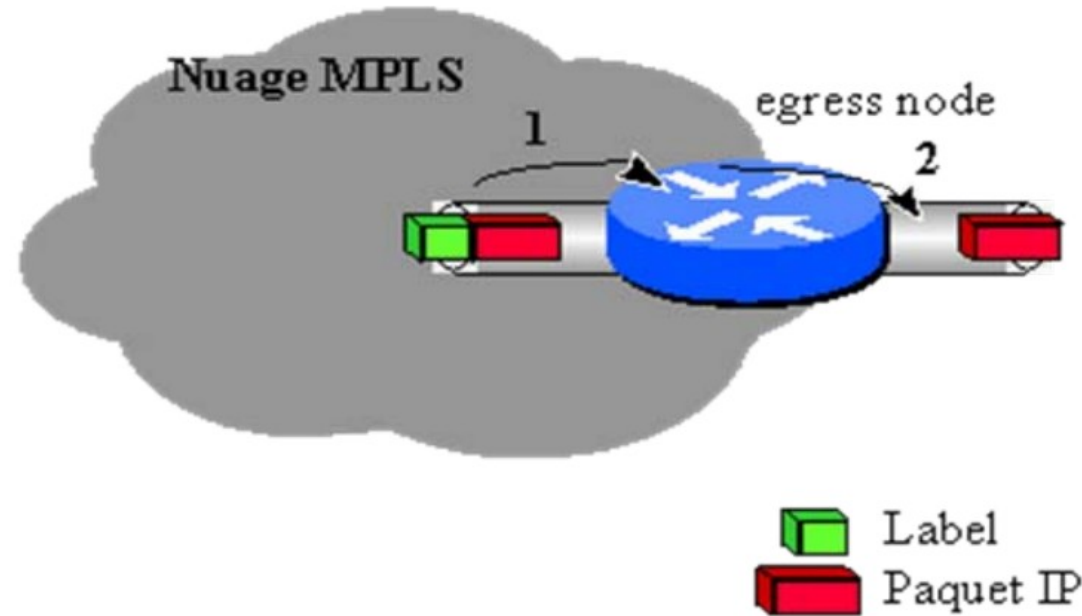
1. Le paquet arrive dans un réseau MPLS
2. En fonction de la FEC auquel appartient le paquet (déterminée à partir de l'adresse IP), l'ingress node consulte sa table de commutation
3. Affecte un label au paquet (3)
4. le transmet au LSR suivant (4).

La commutation de labels sur LSR



1. Le paquet MPLS arrive sur un LSR interne du nuage MPLS
2. le protocole de routage fonctionnant sur cet équipement détermine dans la base de données des labels LIB (*Label Base Information*), le prochain label et prochain LSR ou LER
3. Mise à jour de l'en-tête MPLS
4. Envoie au noeud suivant (LSR ou *l'egress node*)
 - Sur un LSR interne, le protocole de routage de la couche réseau n'est jamais sollicité

Commutation de labels sur LER sortie

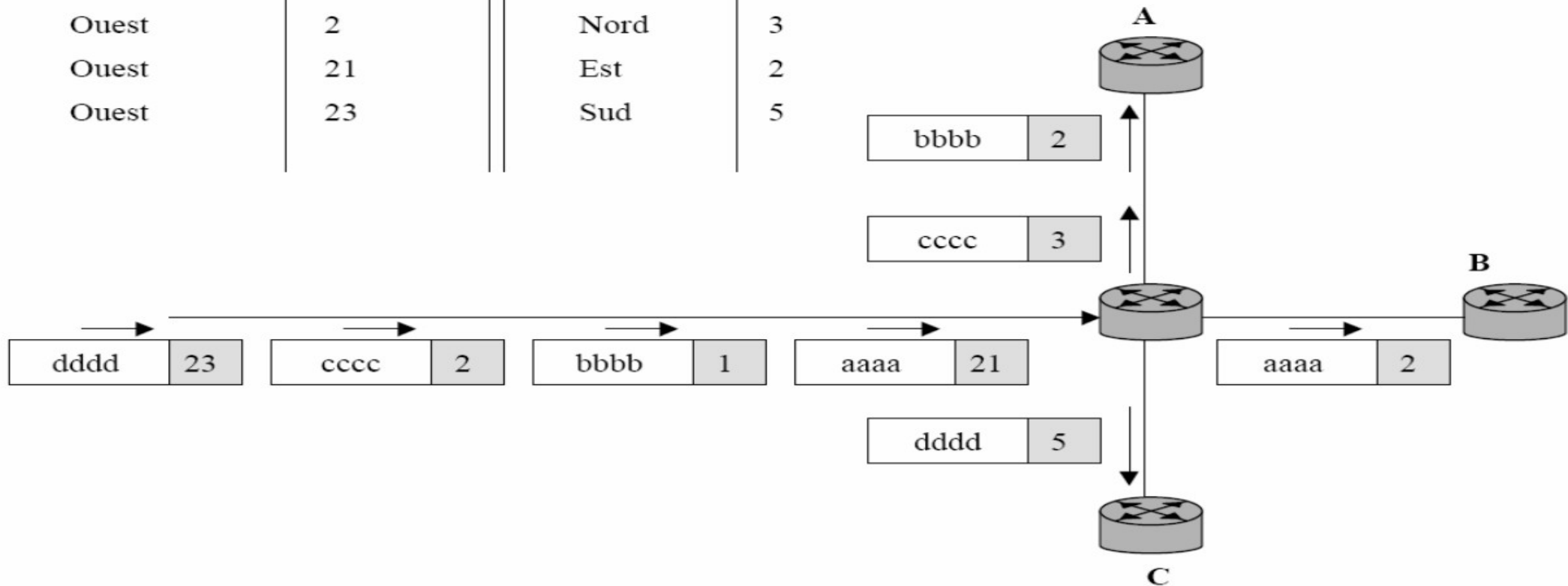


1. Le paquet MPLS arrive à l'*egress node*
2. l'équipement lui retire toute trace MPLS
3. le transmet à la couche réseau

Exemple de commutation

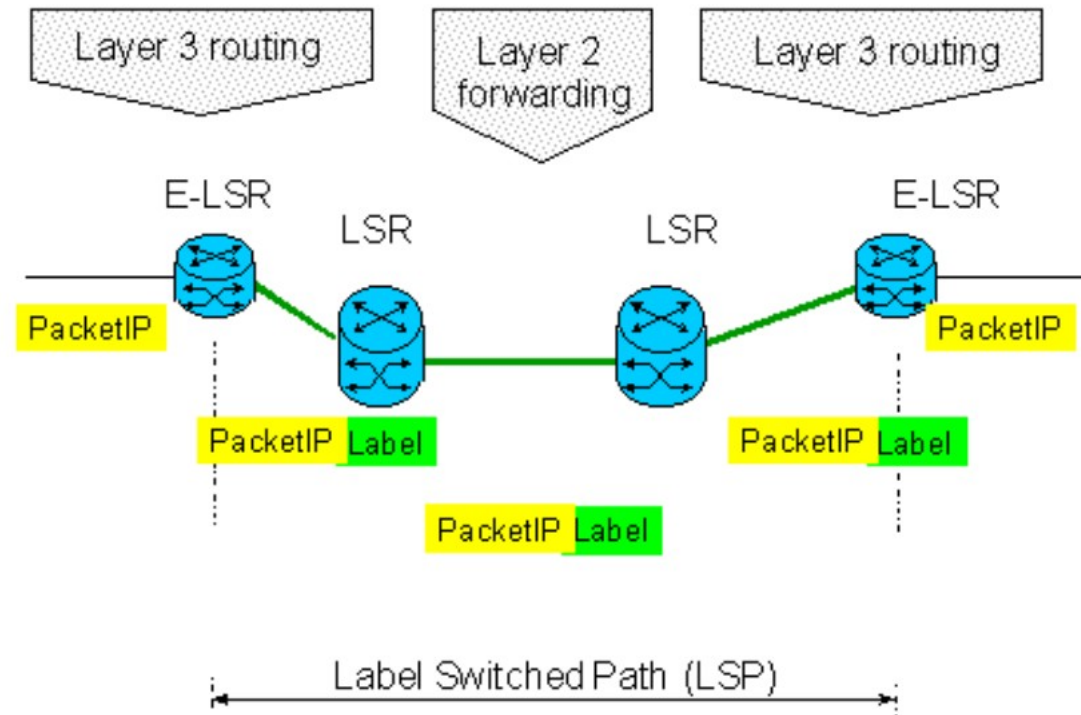
Label Forwarding Table

InPort	InLabel	OutPort	OutLabel
Ouest	1	Nord	2
Ouest	2	Nord	3
Ouest	21	Est	2
Ouest	23	Sud	5



Principes MPLS (2)

- Un flux MPLS est vu comme un flux de niveau 2.5 appartenant niveau 2 et niveau 3 du modèle de l'OSI



Agrégation de flux et routage hiérarchique

■ L'agrégation de flux

- La possibilité de réunir le trafic entrant dans un routeur via plusieurs LSP dans un seul et unique LSP sortant
 - correspond à monter une connexion multi-point à point
 - permet de réduire au maximum le nombre de connexions que les routeurs de coeur de réseau ont à gérer

■ Routage hiérarchique

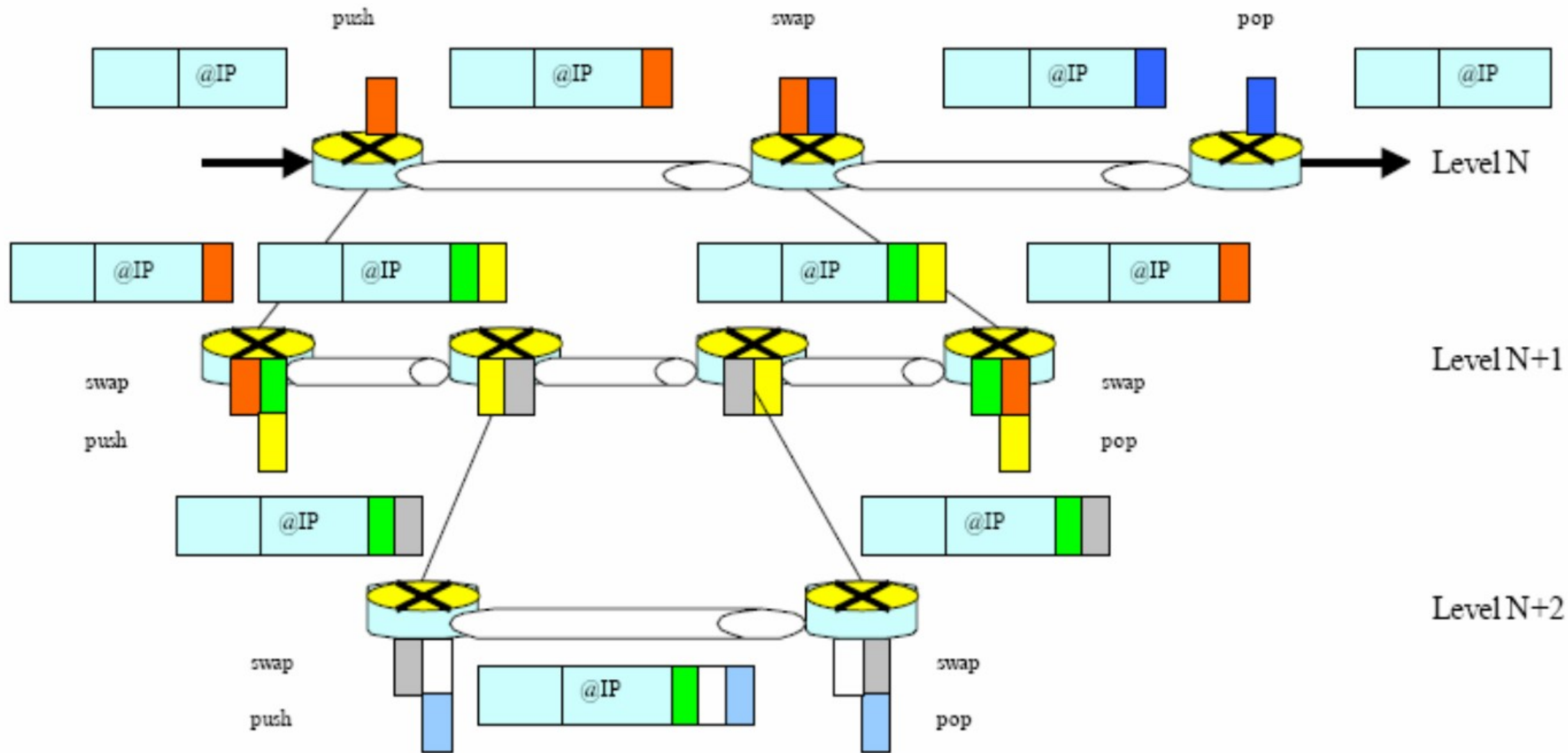
- Consiste à empiler des labels (**champ S, Ethernet**)
 - Permet de construire des LSP, encapsulés dans un autre LSP
 - Permet d'associer plusieurs contrats de service à un flux
- Rappelle ATM : VC dans des VP
 - Cependant, en MPLS, le nombre de niveau d'encapsulation (ou de hiérarchie) n'est a priori pas limité à 2

Commutation hiérarchique

- Empilement d'entête MPLS

Entête de la couche 2	Entête niveau N		
	Entête niveau N+1		
	Entête niveau N+2	Entête de la couche 3	Autres couches, entête + data

Commutation hiérarchique



Découverte du réseau

- En vue de déterminer et de maintenir la connaissance des *ingress* ou *egress nodes (LERs)*
 - utiliser un protocole permettant la découverte du réseau
- MPLS peut s'appuyer sur un ensemble de protocoles déjà disponibles
 - IGP (*interior gateway protocol*)
 - RIP (*routing information protocol*) de type "Vecteur de distance "
 - OSPF (*open shortest path first*) de type "Etat de liens"
- Un ensemble de réseaux MPLS peuvent être interconnectés chacun utilisant leur propre protocole de découverte
 - EGP (*exterior gateway protocol*)
 - BGP (*border gateway protocol*)

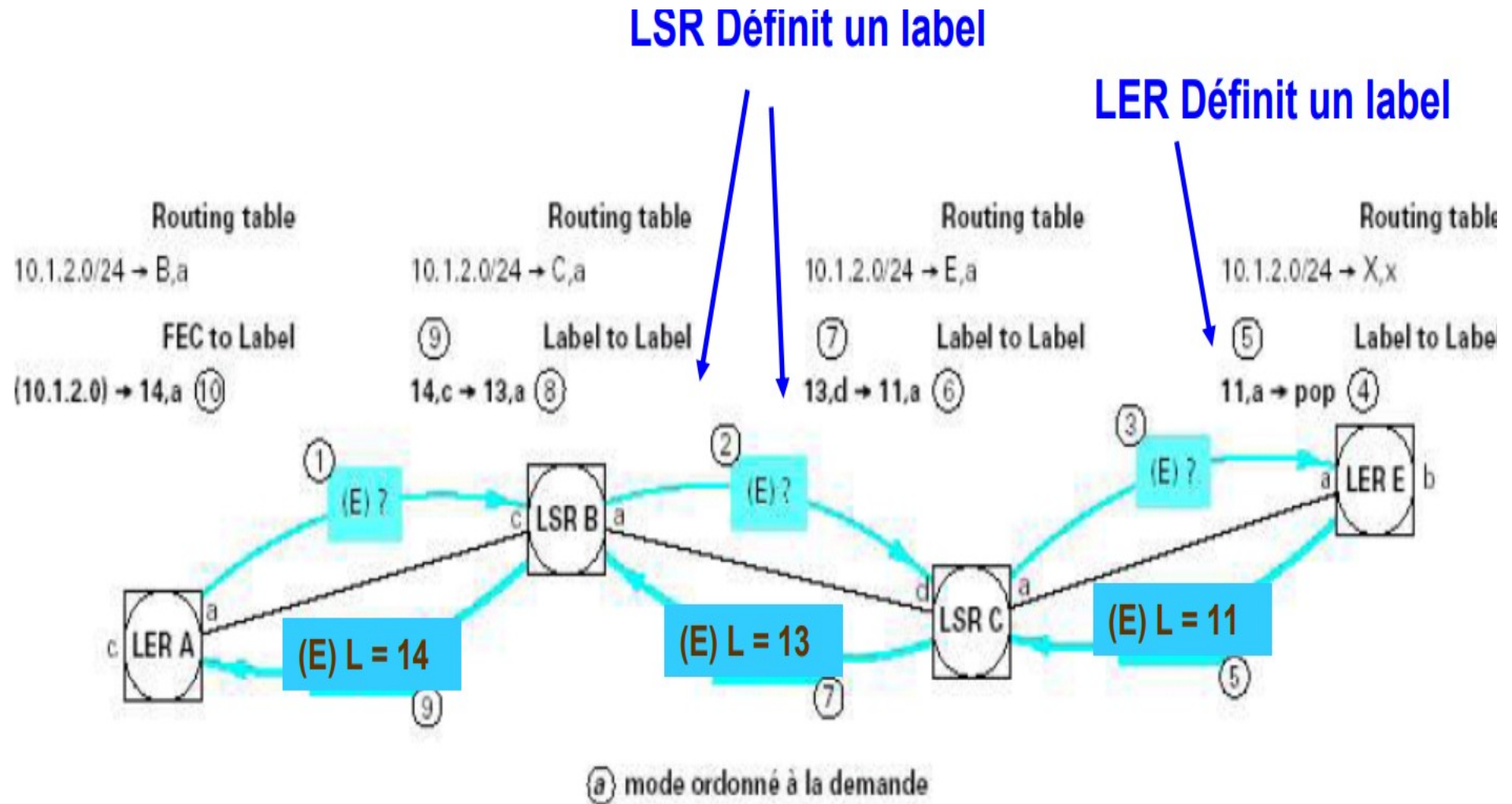
Distribution de Labels (1)

- LDP – *Label Distribution Protocol*
 - Basé sur la table de routage IP pour la distribution de labels
- LDP non ordonné (routage hop-by-hop)
 - Mode indépendant non sollicité
 - Tous le LSR affectent à leurs voisins un nouveau label pour les routes IP découvertes dans leur table de routage IP
 - L'établissement est moins long et moins coûteux en ressource
 - Ne contient pas de paramètres permettant de formuler une demande de ressources
 - L'inconvénient : risque de création de boucles
 - on sait détecter des boucles
 - on ne sait pas les éviter complètement.
- MBGP (MPLS BGP)
 - Collabore avec un autre protocole, souvent le LDP
 - Extension du BGP (*Border Gateway Protocol*)
 - Établit uniquement une association entre deux LER frontières
 - Afin de associer un label commun à une route externe

Distribution de labels (2)

- **CR-LDP** (mode ordonné à la demande)
 - ❑ Établit le chemin bout en bout (LSP)
 1. Un LER découvre dans sa table une adresse dont il n'a pas de label
 2. Demande au prochain LSR (défini dans la table de routage IP) le label correspondant
 3. Ainsi de suite jusqu'à ce que un LSR/LER trouve ou définit un label pour la destination désirée
 - ❑ L'établissement des labels s'effectue alors dans le sens inverse jusqu'au LER qui a effectué la demande
 - ❑ Permet d'associer les caractéristiques de QoS aux chemins

CR-LDP



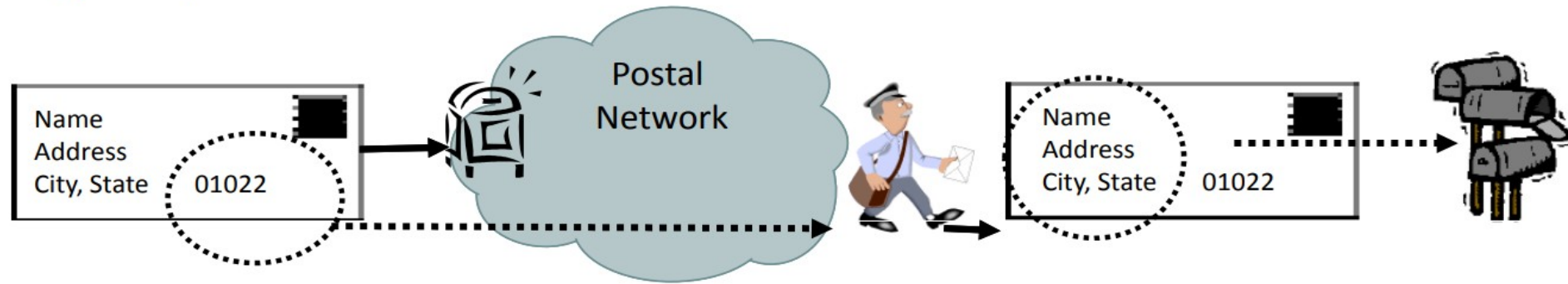
Distribution de labels (3)

■ RSVP-TE

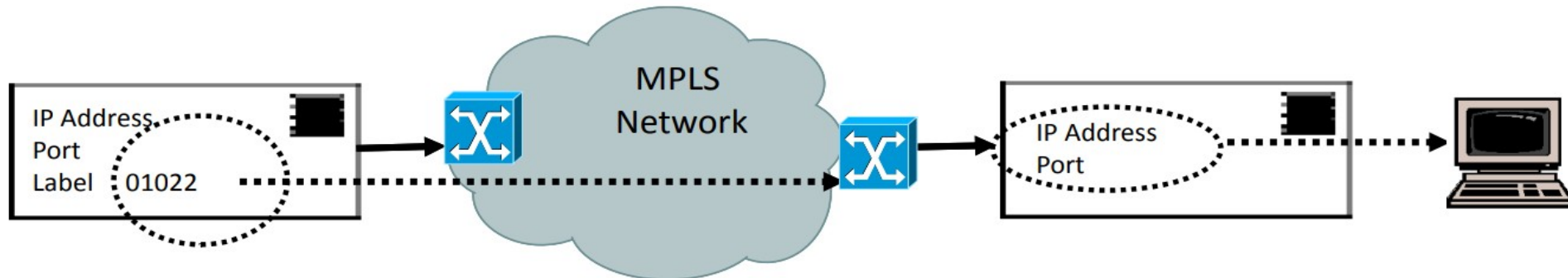
- ❑ Utilise quelques fonctionnalités du protocole LDP
 1. Une demande est faite par le LER d'entrée
 2. Elle traverse le nuage MPLS jusqu'au LER de sortie
 - ❑ utilise les tables de routage
 3. Elle établit la route en inscrivant les information de FEC
 - ❑ Décrit les ressources pour permettre au label déterminé par le LER de sortie de remonter vers le LER d'entrée
- ❑ Permet d'associer les caractéristiques de QoS aux chemins

Efficacité du Switching

Le réseau postal transfère les paquets sur la base du code postal
L'agent de poste se base de son côté sur le nom et l'adresse



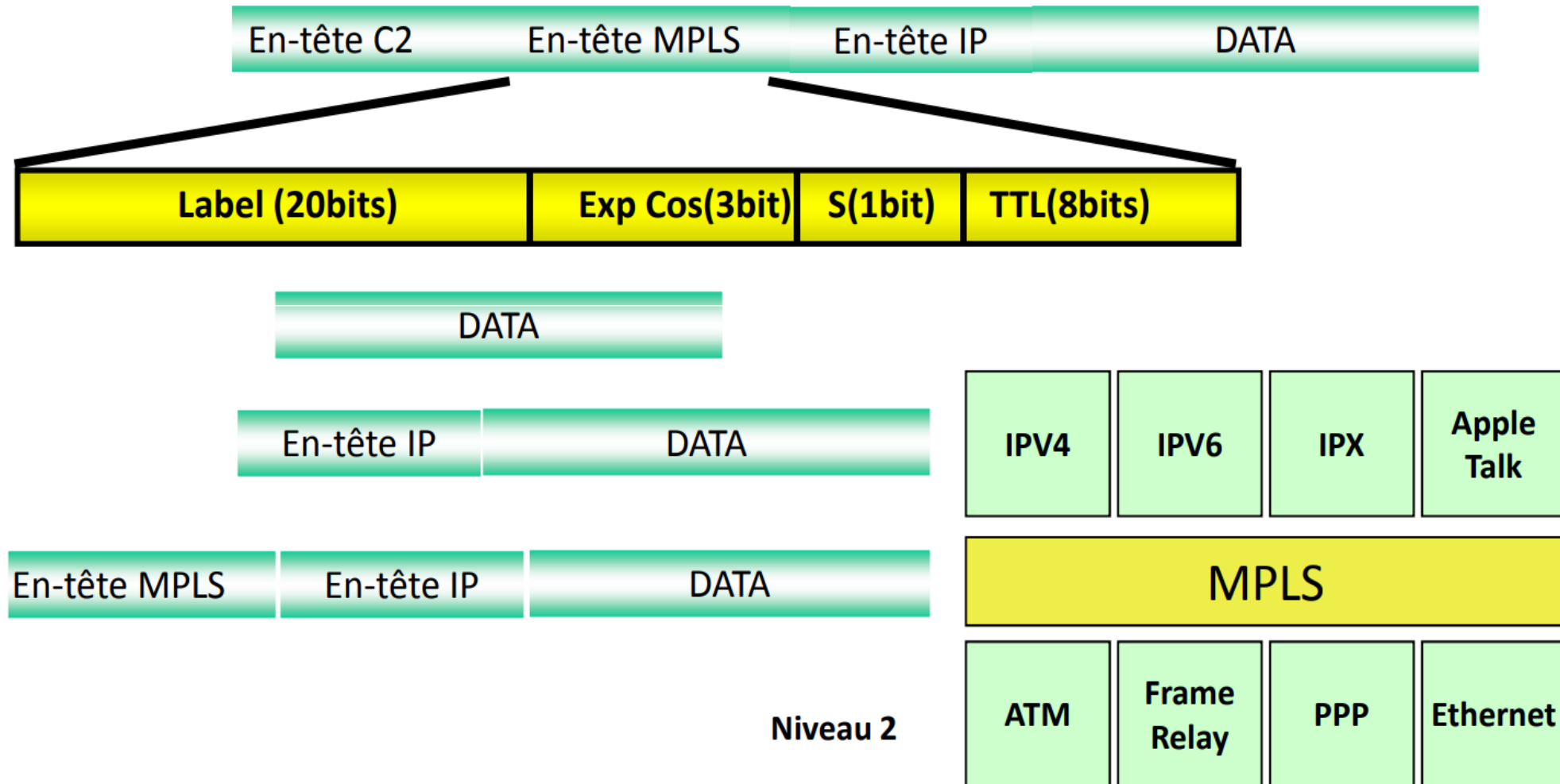
Les switches MPLS ajoutent un label et transfèrent sur la base de la valeur du label
Les Egress switches suppriment les labels et routent sur la base de l'@ IP



Le label MPLS

- ❑ L'architecture MPLS est organisée autour d'une principale notion que constitue **le label ou étiquette**
- ❑ Le format dépend explicitement des caractéristiques du réseau utilisé.
- ❑ Comme les identificateurs VPI/VCI de ATM, le label MPLS n'a qu'une signification locale entre deux équipements MPLS.
- ❑ **L'en-tête MPLS occupe 4 octets (32-bits)**

Le label MPLS



Le label MPLS

- ❑ **Codé sur au moins 32 bits**

- ❑ Avec un ou plusieurs labels codés sur 20 bits chacun

- ❑ *Une PDU MPLS peut contenir une pile de labels*

- ❑ *Le label au sommet de la pile est transmis en premier*

- ❑ » Si $S = 1$, le label codé est au fond de la pile (dernier label)

- ❑ » Si $S = 0$, le label codé est au-dessus d'un autre label

- ❑ **Autres champs**

- ❑ *TTL (Time To Live)*

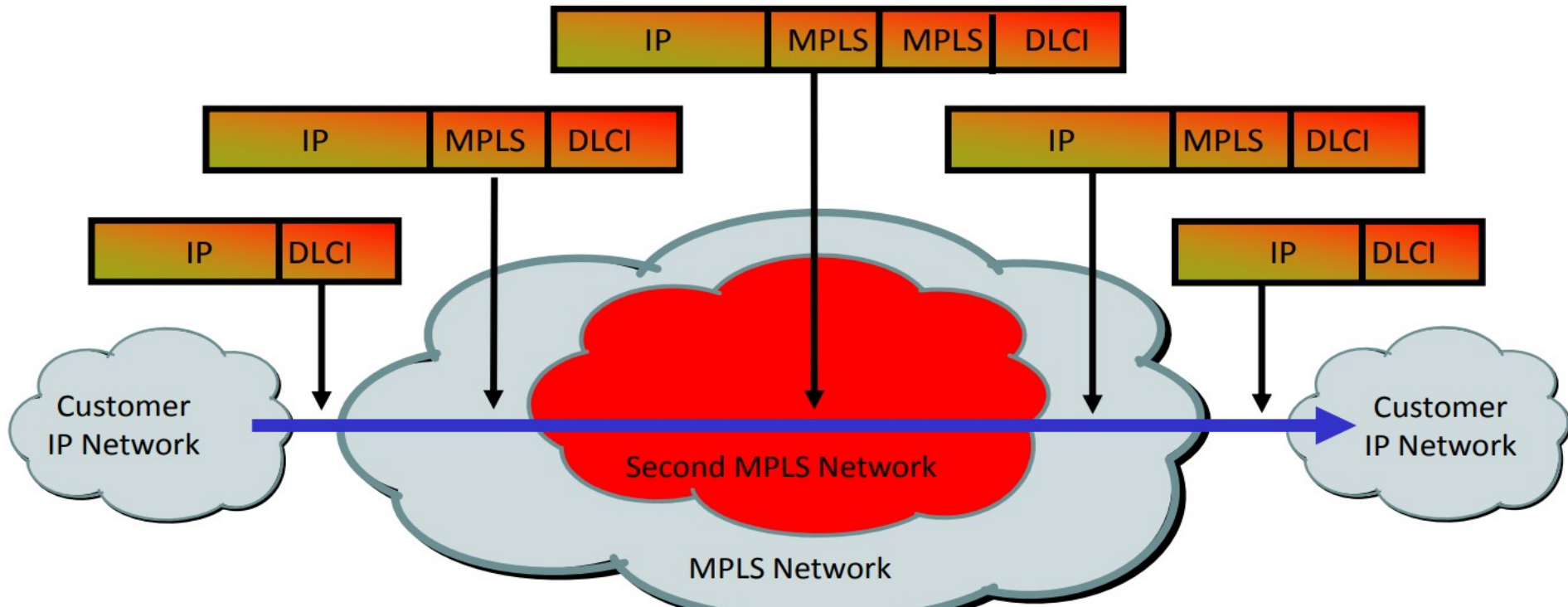
- ❑ Pour éviter que des PDU MPLS puissent boucler indéfiniment dans un domaine

- ❑ *Exp ou COS (Class Of Service)*

- ❑ Pour pouvoir appliquer plusieurs politiques de gestion des files d'attente

Label Stacking

- Label Stacking—processus d'ajout et suppression de labels à chaque fois que le paquet traverse multiple MPLS networks
- Forward est basé sur le label le plus proche de la couche 2

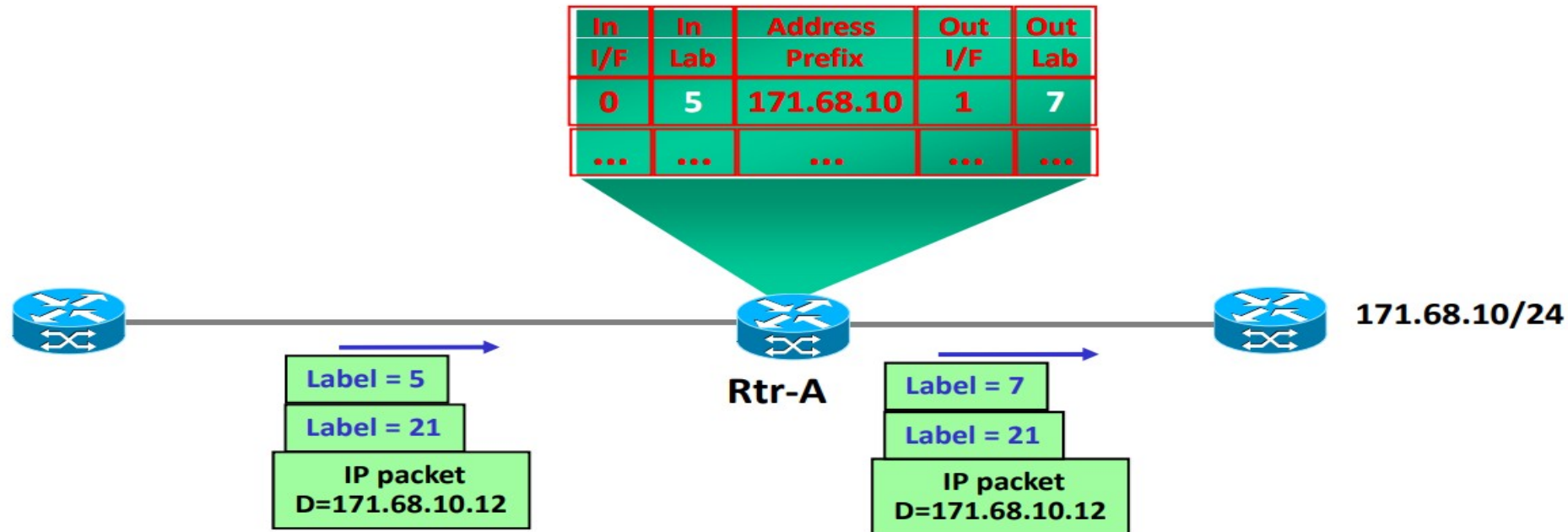


Pile de label



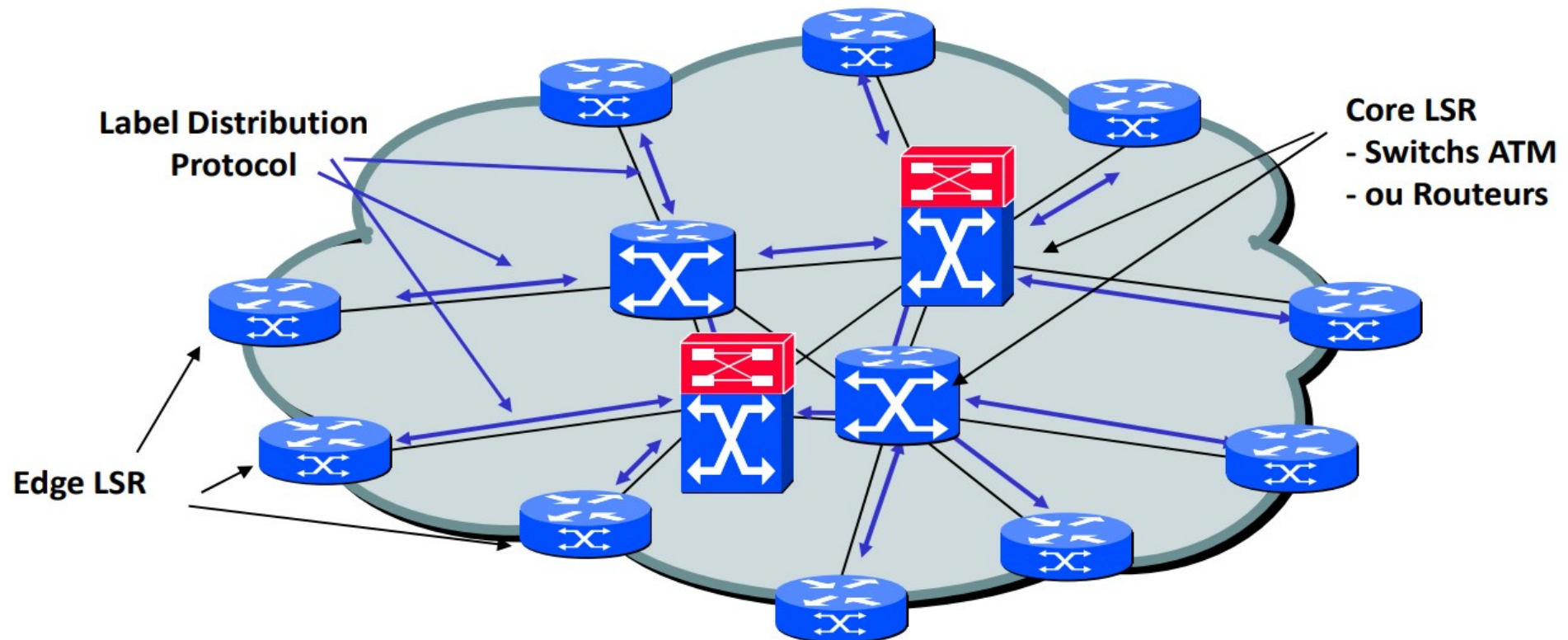
La pile de labels (“label stack”)

- Chaque paquet peut contenir plusieurs niveaux de labels.
 - Le LSR (Rtr-A) commute le paquet labellisé en fonction uniquement du premier label de la pile



Nœuds MPLS

- Edge LSR: Label Switch Router d'extrémité
- Core LSR: Label Switch Router de cœur de réseau



Label Edge Router (LER)

■ Ce sont des routeurs qui traitent le niveau IP et le niveau label

- ❑ Routeur d'extrémité qui joue un rôle important dans l'assignation et la suppression des labels au moment où les paquets entrent dans le réseau ou en sortent.
 1. Décide comment les paquets vont voyager dans le réseau
 2. Reçoit des paquets IP et leur assigne des labels 'Push'
 3. Transfert les paquets labélisés vers le/les routeurs de prochain saut LSR/LER
 4. Reçoit les paquets labélisés à partir des LSR/LER pour les router vers la destination finale 'Pop'
- ❑ Chaque LSR construit une table LFIB (Label Forwarding Information Base).

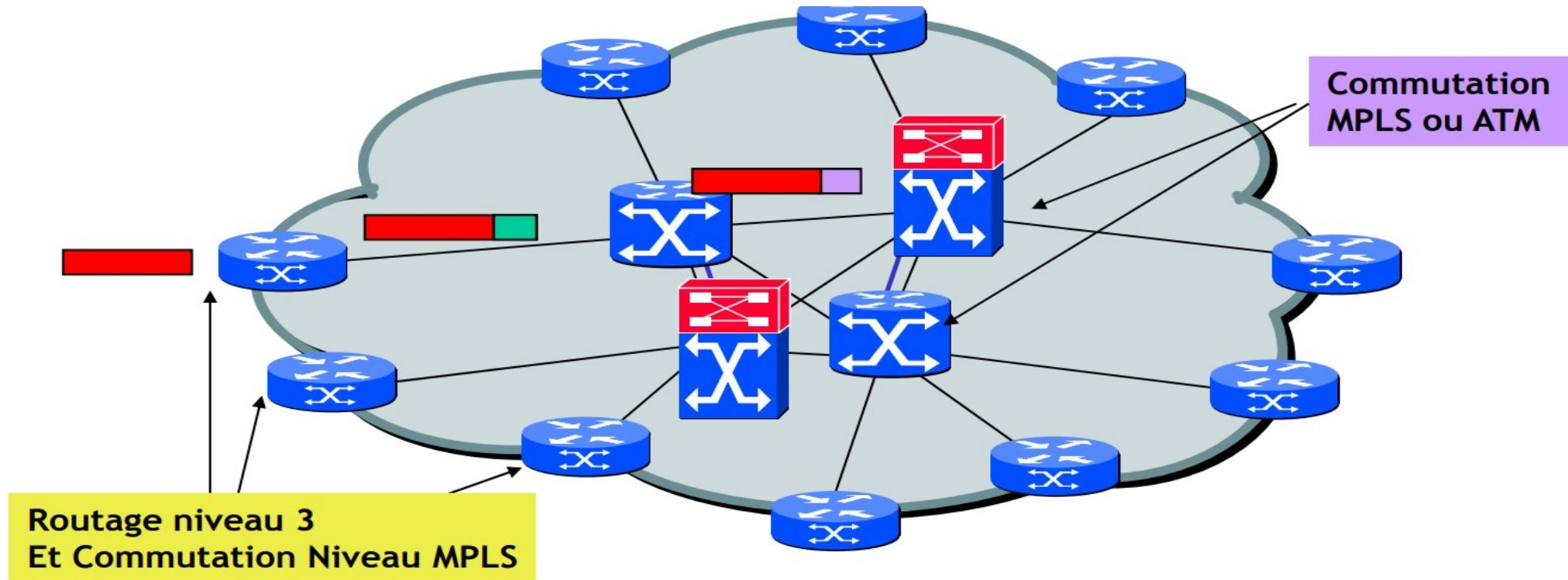


Label Switching Router (LSR)

- ❑ Routeur ou commutateur capable de **commuter des paquets** ou des cellules, en fonction de la valeur des labels qu'ils contiennent.
- ❑ Dans le cœur du réseau, les LSRs procèdent tout simplement à **la lecture/écriture et la commutation des labels**, et non les adresses des protocoles de niveau supérieur.
- ❑ Chaque LSR construit une table LFIB (Label Forwarding Information Base).

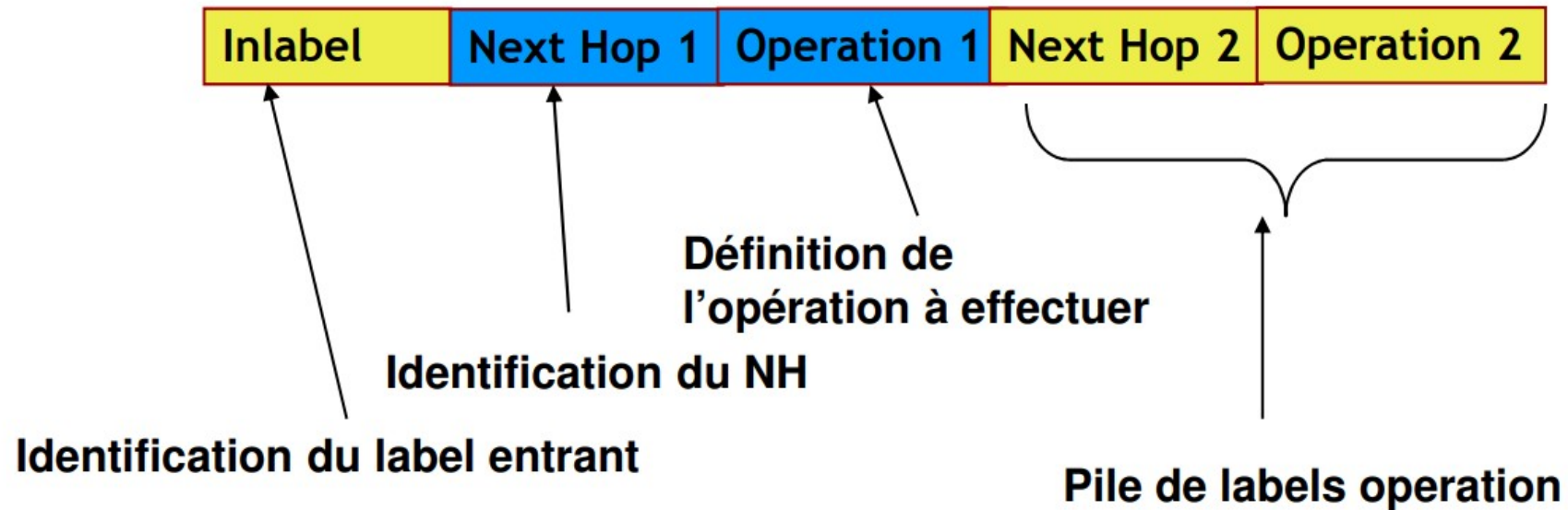


LSR et LER : Commutation et routage

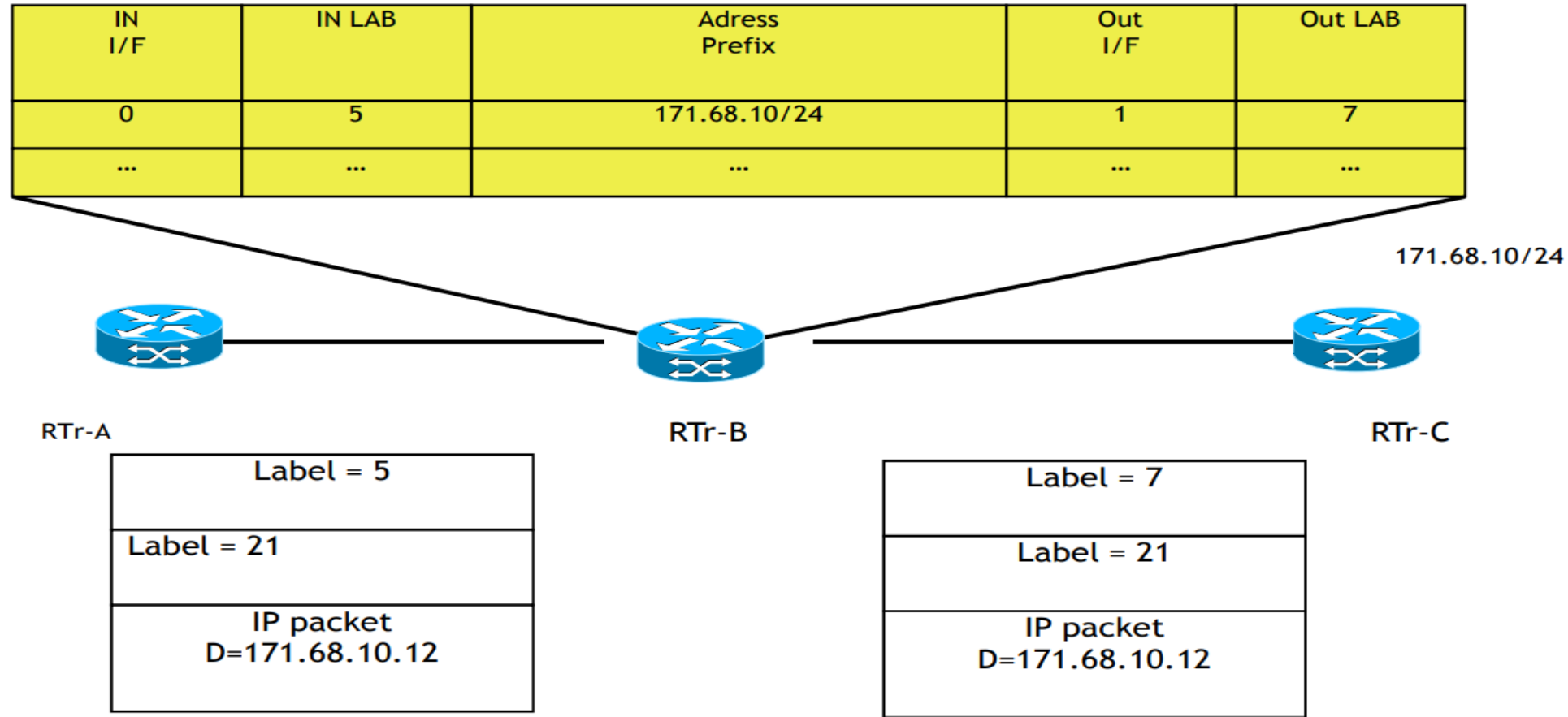


- ❑ **Avantages**
 - ❑ Calcul unique au niveau de l'entrée du réseau
 - ❑ Rapidité dans le cœur de réseau
- ❑ L'intelligence se trouve aux extrémités du réseau

Label swapping

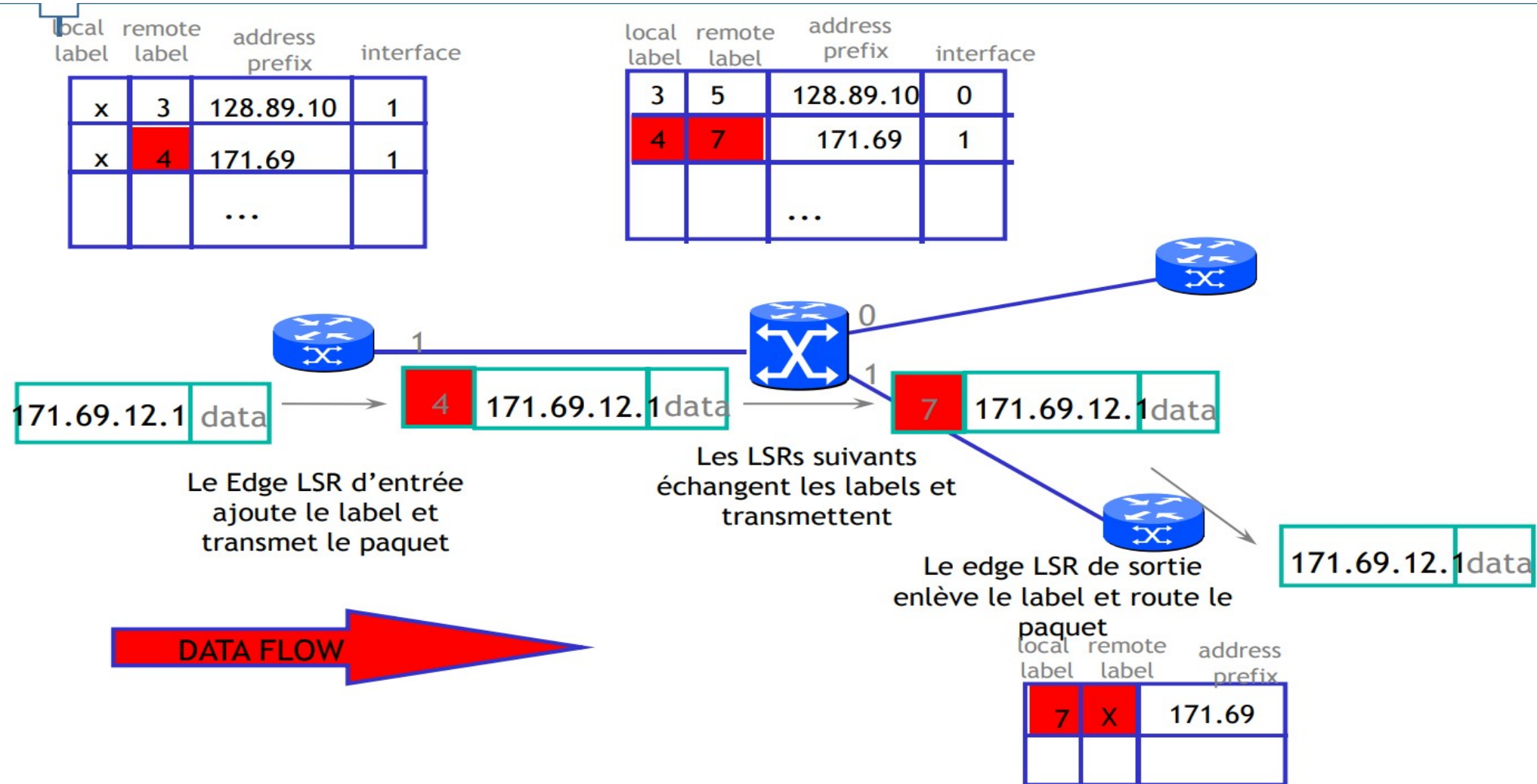


Label Swapping



**Processus de transfert des paquets vers leur destination
sur la base de la lecture et écriture des labels**

La commutation des paquets



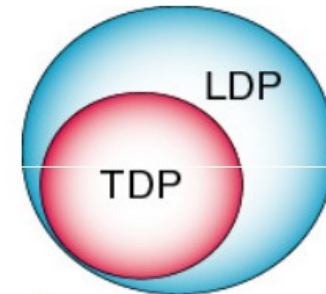
Distribution de labels

❑ Problème posé

- ❑ Comment coordonner les tables de commutation des LSR ?
 - ❑ *Par distribution d'associations entre FEC découvertes et labels (en anglais, FEC-Label mappings)*

❑ Protocoles proposés

- ❑ TDP/LDP (Tag/Label Distribution Protocol)
 - ❑ *TDP est un protocole CISCO propriétaire*
 - ❑ *LDP a été défini par l'IETF dans la RFC3036*
 - ❑ *Utilisés notamment dans le cas de FEC définies par sous-réseau*
- ❑ Extensions de RSVP (Resource Reservation Protocol)
 - ❑ *Permet d'établir des LSP en fonction de critères de ressources et de qualité de service*
- ❑ Extensions de BGP (Border Gateway Protocol)
 - ❑ *Permet de distribuer les labels en même temps que les routes découvertes, mais aussi l'échange de routes VPN*



Les protocoles de distribution de label

- ❑ LDP

- ❑ associe des labels à des FECs (adresses unicast)

- ❑ RSVP

- ❑ Utilisé pour le “Traffic Engineering” et la réservation de ressources

- ❑ PIM

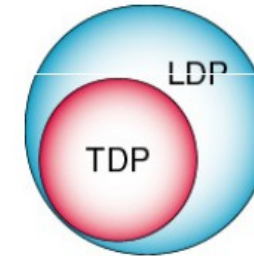
- ❑ association de label en multicast

- ❑ BGP

- ❑ Labels externes: VPN

La distribution de labels

- ❑ La distribution des labels aux LSR est réalisée grâce au protocole LDP
 - ❑ Basé sur le protocole TDP (Tag Distribution Protocol) de Cisco RFC 3036
 - ❑ *TDP et LDP diffèrent principalement par le format des paquets ou messages qu'ils transmettent*



- ❑ **Fonctions principales**

- ❑ **Découverte de voisins**

- ❑ *Pour un LSR donné, déterminer si ses voisins directs sont également des LSR ou seulement des routeurs classiques*

- ❑ **Distribution de labels**

La distribution de labels

- ❑ LDP définit une suite de procédures et de messages utilisés par les LSR pour s'informer mutuellement du mapping entre les labels et le flux.
- ❑ Les labels sont spécifiés selon le chemin " Hop By Hop " défini par l'IGP dans le réseau.
- ❑ Chaque nœud doit donc mettre en œuvre un protocole de routage de niveau 3, et les décisions de routage sont prises indépendamment les unes des autres.

Label Description Protocol (LDP)

RFC3036: Principes

- ❑ Le protocole LDP est utilisé pour associer un FEC (Fowarding Equivalence class) à un label.
- ❑ Les labels permettent de créer des chemins labellisés LSPs
- ❑ LDP fonctionne en mode session entre “peers” .
 - ❑ Ces peers peuvent ne pas être adjacents.
- ❑ LDP est bi-directionnel et permet la découverte dynamique des noeuds adjacents grâce à des messages Hello échangés par UDP.
- ❑ Une fois que les 2 nœuds se sont découverts, ils établissent une session TCP
- ❑ Cette session agit comme un mécanisme de transport fiable :
 - ❑ des messages d'établissement de session TCP,
 - ❑ des messages d'annonce de labels
 - ❑ et des messages de notification.

RFC3036: Les messages LDP

- ❑ Les peers échangent des messages LDP : 4 types de messages
- 1. **Message de découverte (Discovery):** annonce et maintient une adjacence entre des LSRs du réseau MPLS (notion de LDP peers)
 - ❑ **Mécanisme de base:** localisation de LSRs sur le même lien pour établir une adjacence
 - ❑ Un paquet “link Hello (en UDP)” est envoyé **à l’adresse multicast** “all-routers-in-subnet”
 - ❑ La réception d’un link Hello identifie un peer LDP potentiel: Contient le “LDP identifier” et le “label space”
 - ❑ **Mécanisme étendue:** localisation de LSRs distants
 - ❑ Un paquet “targeted Hello (en UDP)” est envoyé à une **adresse spécifique** sur le port UDP: LDP discovery (port 646)
 - ❑ Le “targeted LSR” répond en envoyant des “targeted Hellos” à l’émetteur

Les messages LDP : Découverte

❑ Principe de base

1. Tout LSR transmet périodiquement à ses voisins
 - ❑ Un message **LDP Hello** avec comme adresse destination une adresse « multicast » réservée
2. Tout LSR recevant un message LDP Hello
 - ❑ Répond par un autre message LDP **link Hello** s'il utilise lui-même LDP
- ❑ Etant donné deux LSR voisins utilisant LDP
 - ❑ Le LSR avec la plus grande adresse IP
 - ❑ Devient actif et établit une connexion TCP sur le port 646 (711 pour TDP) pour la transmission de messages LDP
 - ❑ Le LSR avec la plus petite adresse IP
 - ❑ Devient passif et attend l'établissement de la connexion TCP avec son voisin actif
 - ❑ **Note** : des options peuvent être négociées pendant l'établissement de la connexion entre LSR voisins
 - ❑ Modes de distributions des labels, valeur du temporisateur « keep-alive », etc.

Les messages LDP: (suite)

2. Message de mise en Session (TCP)

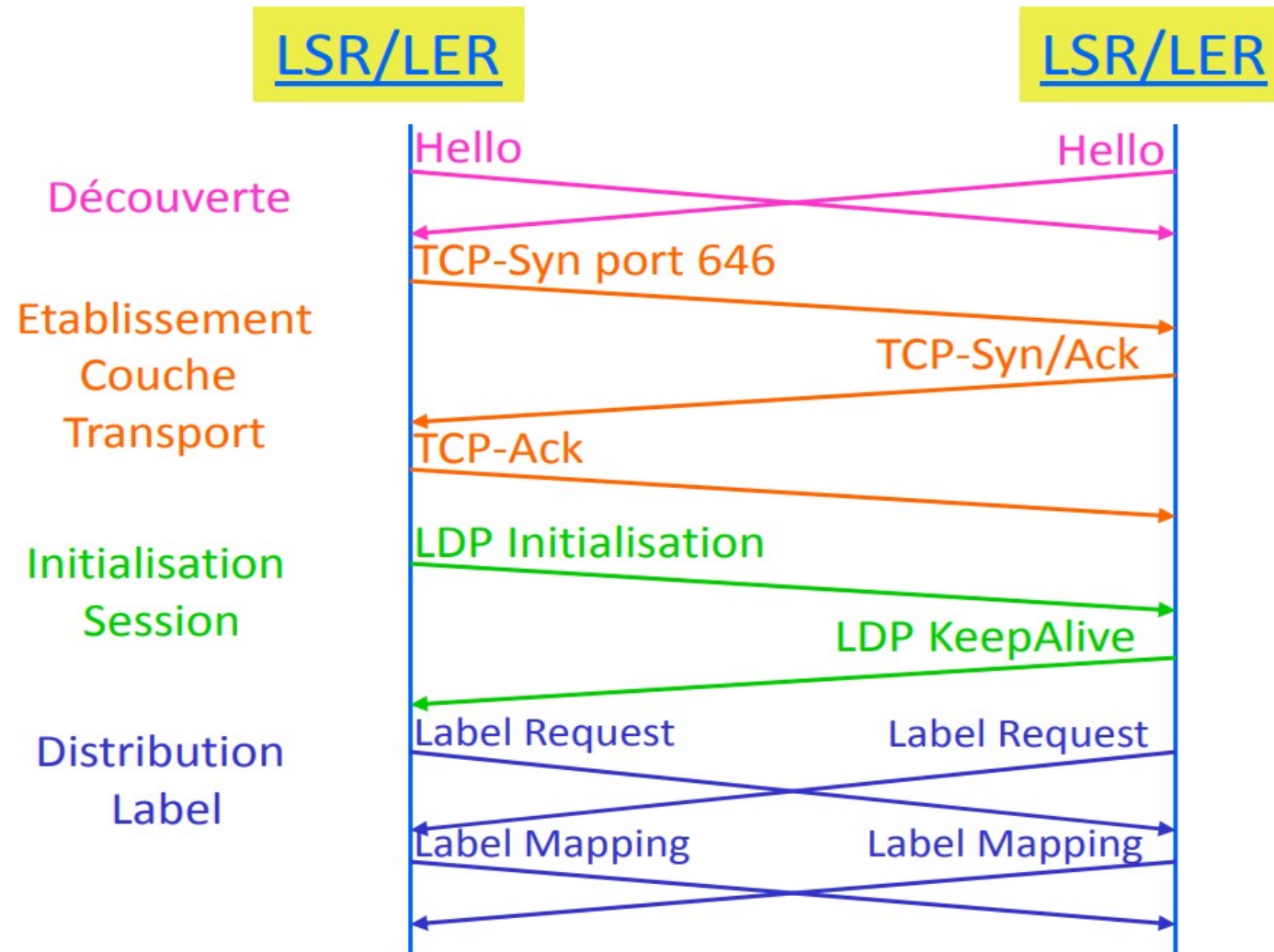
- ❑ Consécutif à la phase “discovery” : découverte d’une adjacence entre LDP Peers
- ❑ Etablit, maintient et termine les sessions entre LDP peers
- ❑ Mise en session en Deux phases:
 - ❑ Etablissement d’une connexion TCP (port 646)
 - ❑ Initiatilisation de la session: négociation des paramètres de sessions: Version de protocole LDP, méthode de distribution de label, label range, ...

3. Messages d’annonces (“Notification message”)

- ❑ Permet de créer, modifier et supprimer l’association entre label et FEC

4. Messages de notification d’erreur

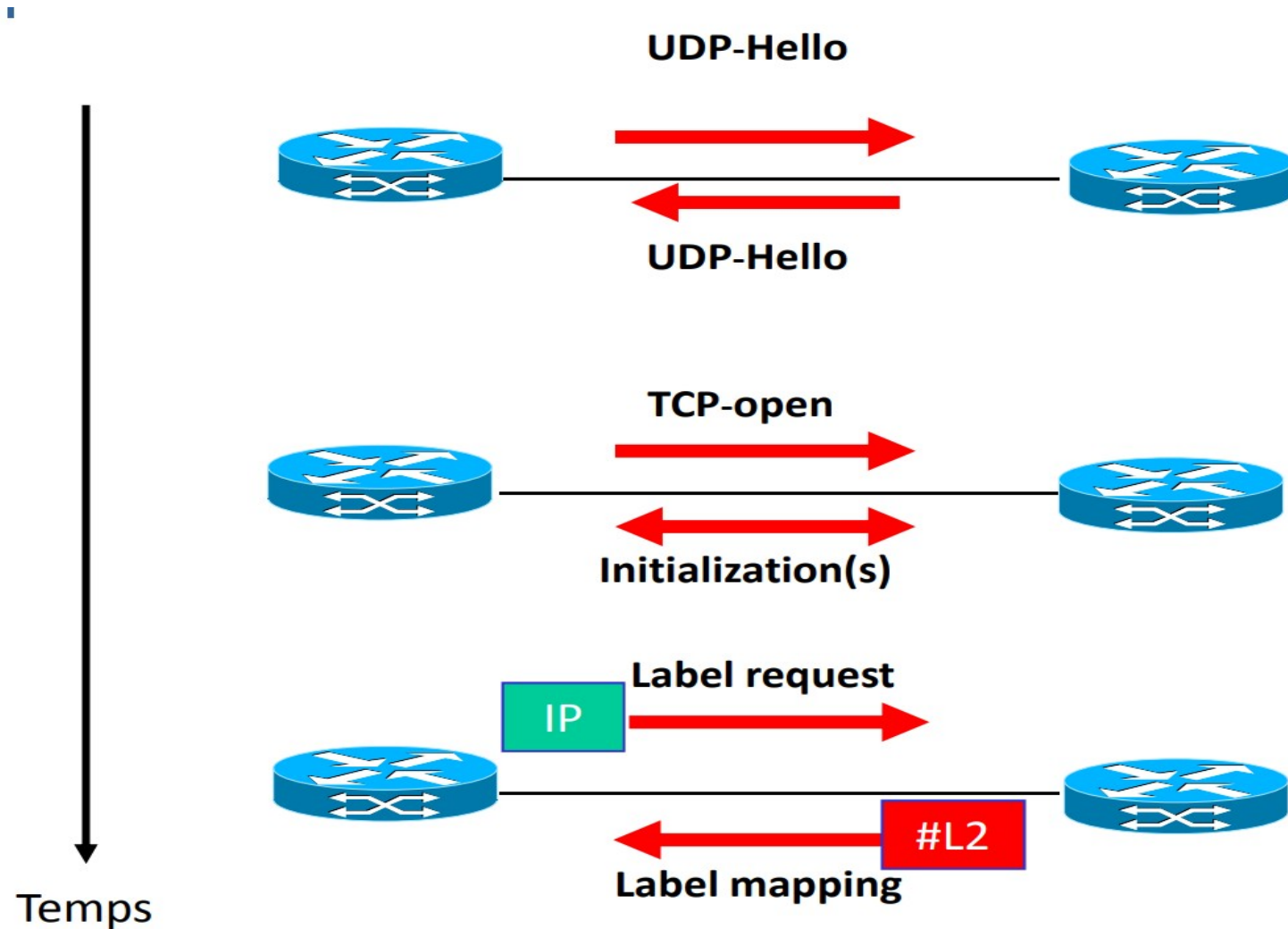
LDP Initialization



TCP :

- fiabilité
- contrôle de flux

Les messages LDP: (suite)



Messages LDP: suite

- ❑ Principaux types de messages

- ❑ Message « Hello »

- ❑ *Utilisés pour la découverte de voisins*

- ❑ Message « Keep-alive »

- ❑ *Transmis périodiquement en l'absence d'autres messages (panne)*

- ❑ Message « Label Mapping »

- ❑ *Utilisé par un LSR pour annoncer une association FEC-Label*

- ❑ Message « Label Withdrawal »

- ❑ *Utilisé par un LSR pour retirer une association annoncée auparavant*

- ❑ Message « Label Release »

- ❑ *Utilisé par un LSR pour indiquer qu'il n'utilisera plus une association reçue précédemment*

- ❑ Message « Label Request »

- ❑ *Utilisé par un LSR pour demander un label pour un FEC donné*

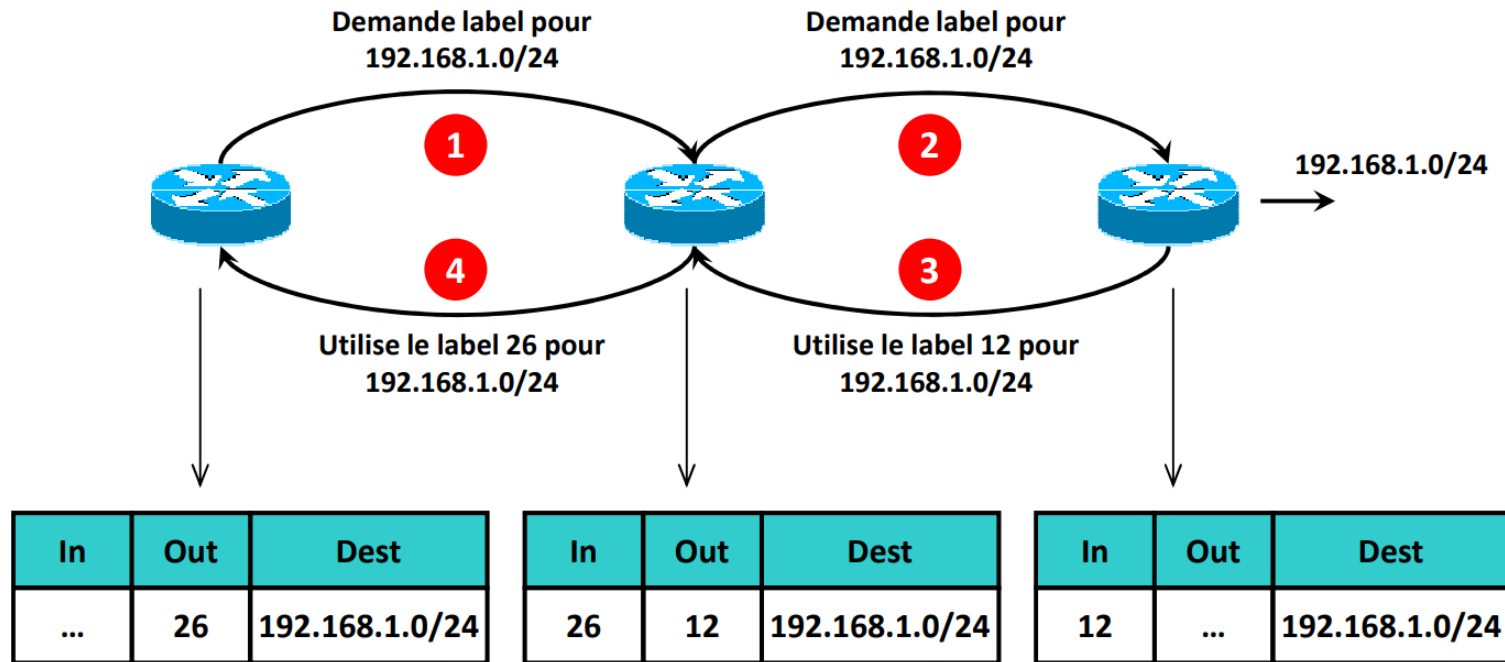
Les sessions LDP entre nœuds distants

- ❑ Les sessions LDP peuvent être établies entre des "peers" qui ne sont pas directement connectés.
- ❑ Peut être utilisé par le Traffic Engineering
- ❑ Le mécanisme de découverte est différent
 - ❑ message "hello" pour des nœuds connectés directement
 - ❑ message "targeted hello" pour des nœuds distants

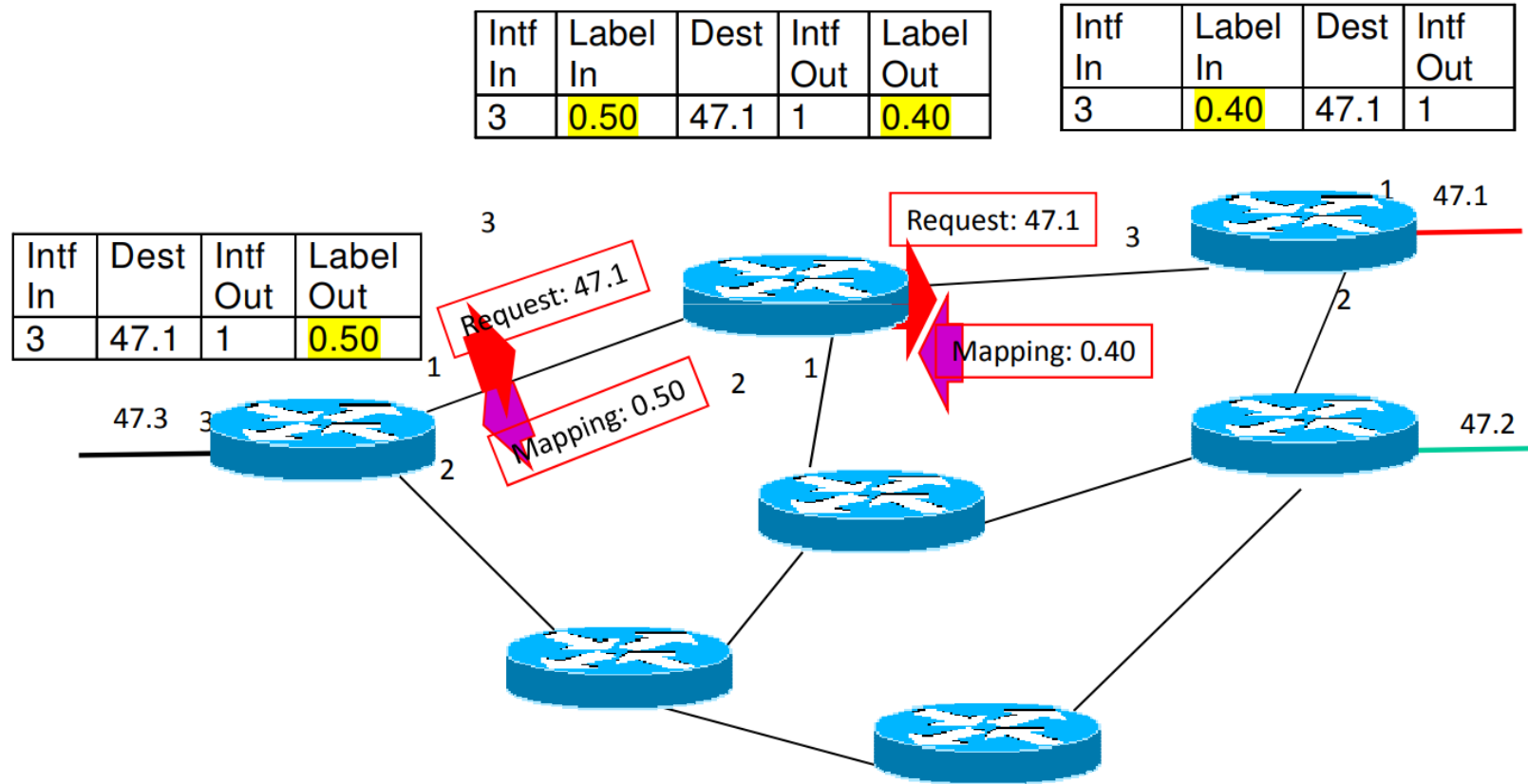
Procédures de distribution des labels

I

Downstream on-demand



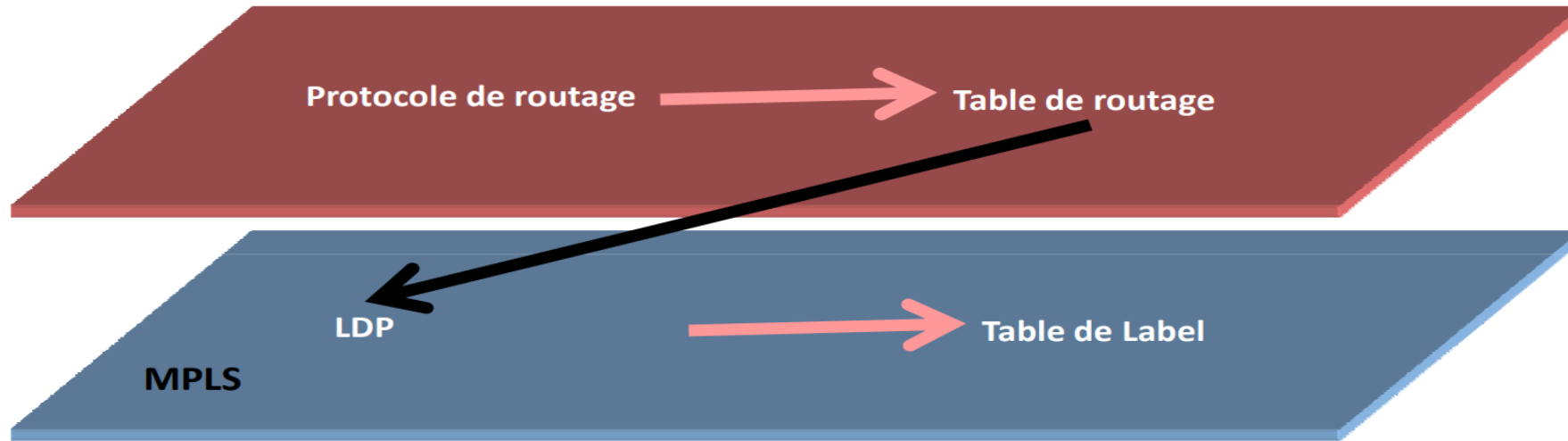
MPLS Label Distribution



LDP Identifiers et adresses de Next-Hop

- ❑ Un LSR enregistre les labels recus dans la “Label Information Base (LIB)”
 - ❑ (LDP Identifier, Label), IP prefix
- ❑ Le LSR associe l’adresse du “next-hop” avec l’identifiant du LDP Peer afin de retrouver le label associé dans la LIB.
- ❑ Afin de permettre cette association, les LSRs annoncent leur adresses d’interface via LDP.

Récap MPLS



1. Protocole de routage
2. Construction de la RIB -> @dest/Next-Hop/Interface
3. Protocole LDP : association Next Hop, label
4. Construction de la RIB -> @dest/Next-Hop/Label/Interface

Fonctionnalités avancées du MPLS (1)

- Le choix de MPLS par les opérateurs est motivé principalement par :
 - ❑ Les possibilités de gestion du trafic ou d'Ingénierie de Trafic
 - Privilégie certains aspects du réseau selon le contexte
 - ❑ éviter les points de forte congestion en répartissant le trafic sur l'ensemble du réseau
 - ❑ utiliser efficacement des ressources du réseau
 - ❑ privilégier certaines routes moins cher
 - ❑ La mise en œuvre de la Qualité de Service
 - permettre d'associer des ressources et de garantir de la QoS sur un LSP

Fonctionnalités avancées du MPLS (2)

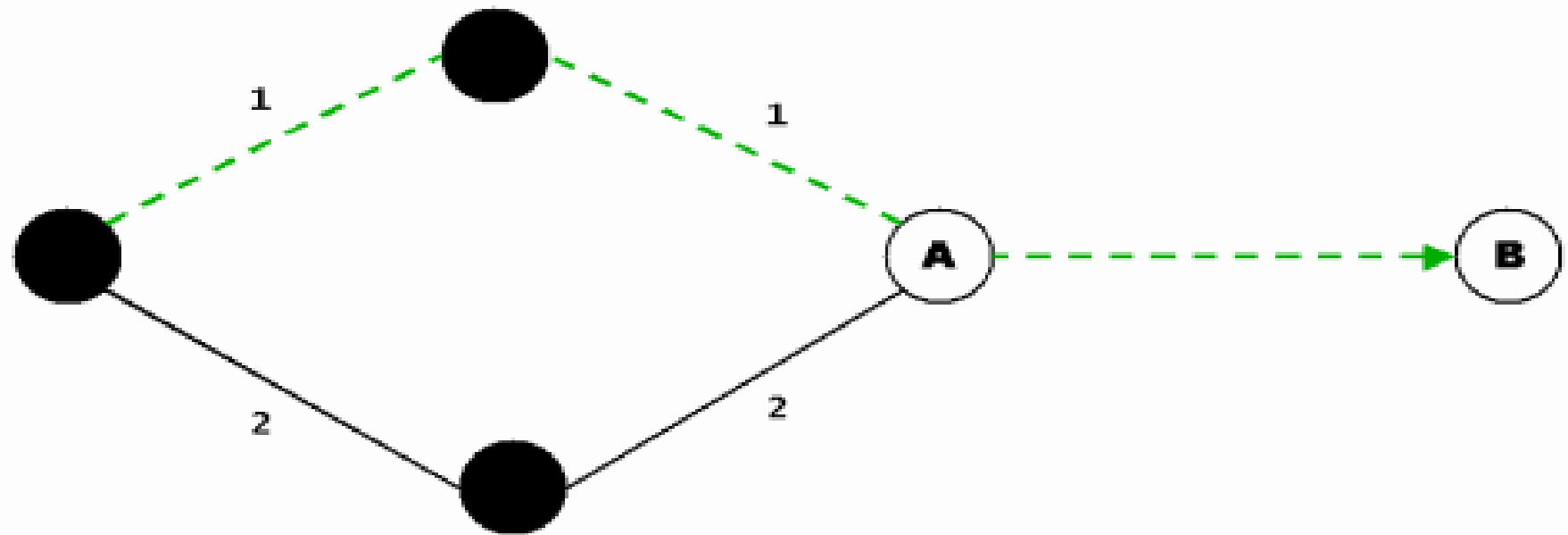
- Ingénierie de trafic

- Un peu d'histoire de *Traffic Engineering*

- Réalisé grâce à des métriques de liens associées à des protocoles de routage internes (RIP, OSPF,...)
 - Fin des années 90, possibilité de le faire avec des technologies de niveau 2 : ATM et Frame Relay
 - Aujourd'hui, MPLS est un nouveau mécanisme de Traffic Engineering
 - Offre une plus grande flexibilité de routage IP (bande passante, QoS,...)

Ingénierie du trafic

- **Ingénierie de trafic**
 - **But : optimiser l'utilisation des ressources du réseau**



Fonctionnalités avancées du MPLS (3)

- Qualité de Service

- Deux approches ont été retenues à l'IETF

- CR-LDP : *Constraint based Routing* LDP, définit des extensions à LDP

- LDP + *Traffic Engineering*

- RSVP-Tunnels, définit des extensions à RSVP pour la commande de LSP

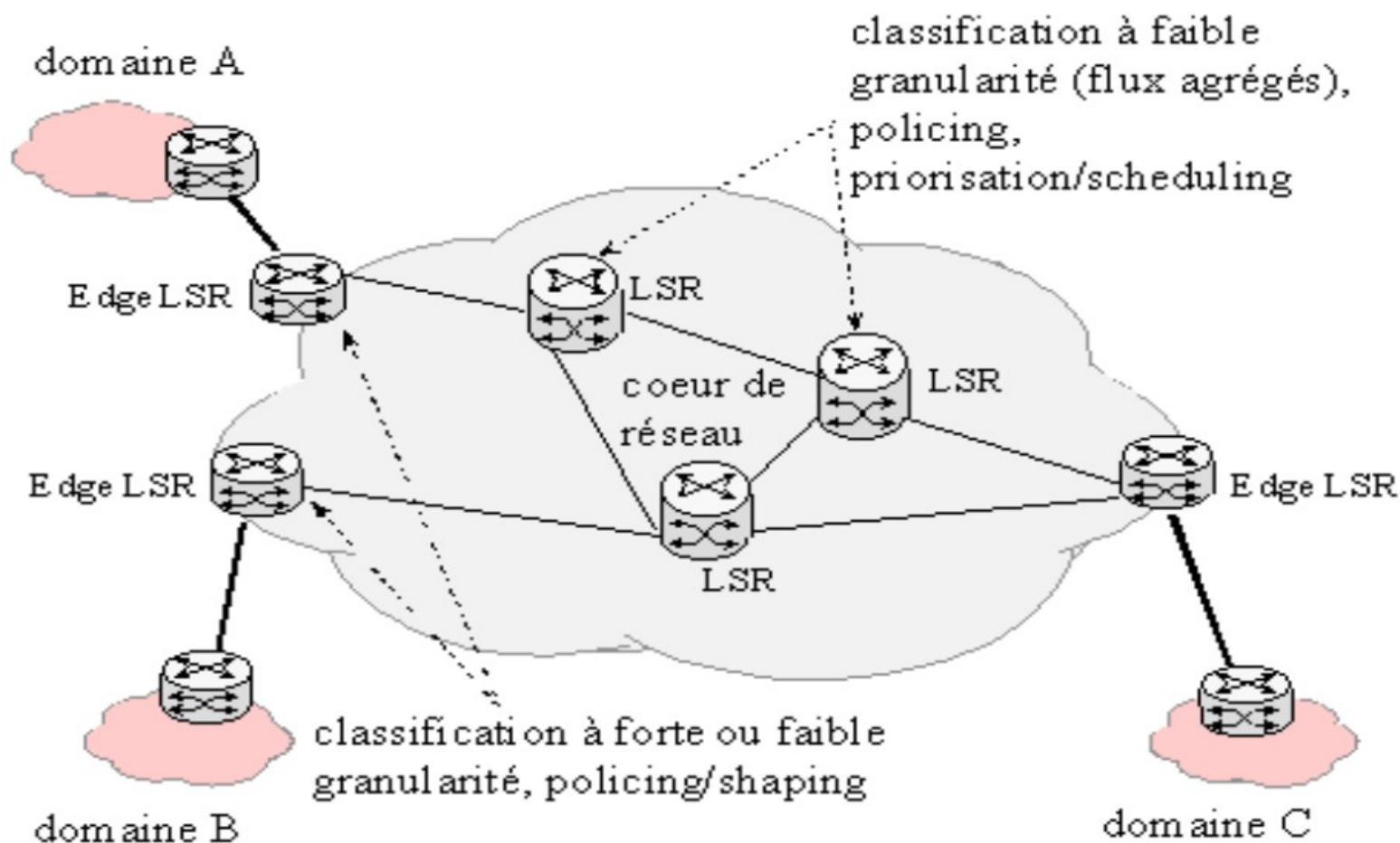
- DiffServ

- Mécanismes complémentaires permettant d'établir une QoS consistante sur les réseaux MPLS

MPLS DiffServ

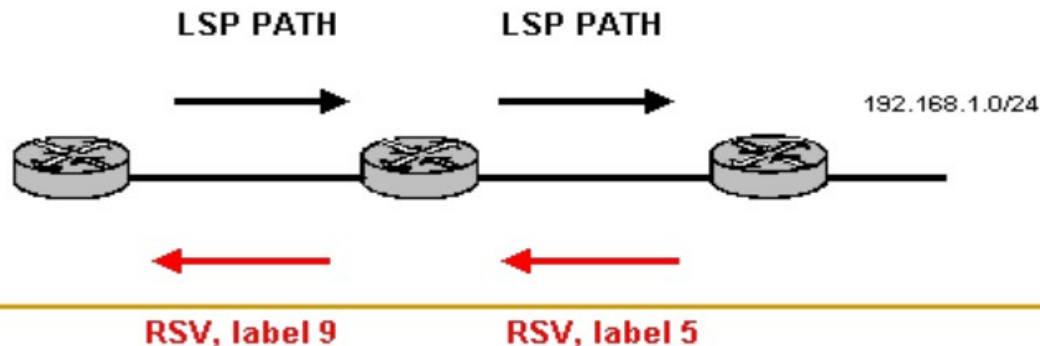
- MPLS est amené à inter fonctionner avec DiffServ
 - car LDP supporte avant tout de la QoS à faible granularité

INTEGRATION MPLS / DIFFSERV



MPLS et RSVP (IntSev)

- Extensions des protocoles pour traiter la QoS : RSVP-TE
 - Modifié pour supporter la gestion de QoS au sein d'un réseau MPLS
 - La source réclame un chemin vers le destinataire (chemin LSP PATH)
 - Message est reçu par le destinataire :
 - envoie d'un message de réservation sur le chemin LSP (LSP RESV) vers la source en réservant la bande passante nécessaire sur les liens intermédiaires
 - Message reçu par la source :
 - le flot de données peut être transféré



VPN : Les offres

- Répondent aux besoins des entreprises nécessitant :
 - ❑ Un service de réseau IP privé sur une infrastructure de réseau IP public
 - ❑ Un passage à l'échelle aisé
 - ❑ QoS
 - ❑ Accès contrôlé vis à vis des autres flux sur l'infrastructure public

VPN s'appuyant sur le protocole MPLS

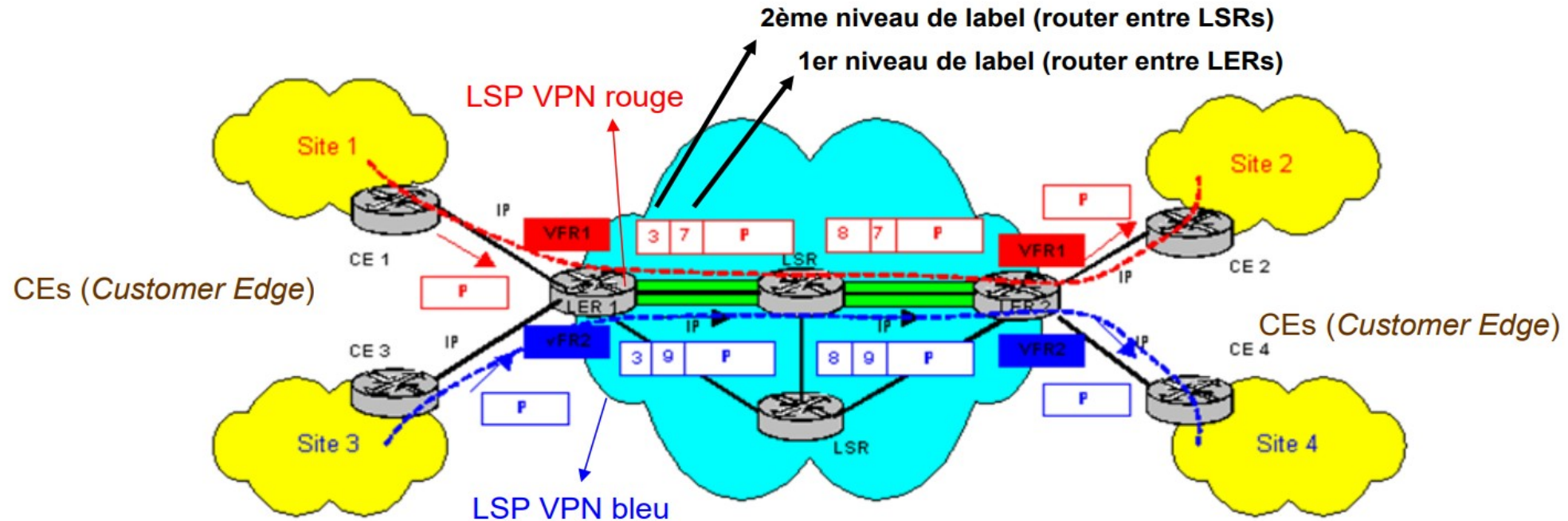
- Différents de l'image classique du VPN
 - S'appuient sur la séparation des paquets sur la valeur de leurs labels
 - Louée que par les LSRs qui appartiennent à un LSP précis
- MPLS-VPN permet d'isoler le trafic entre sites n'appartenant pas au même VPN
 - Permet d'utiliser un adressage privé et d'avoir un recouvrement d'adresses entre différents VPNs.
- Types de VPN MPLS :
 - les VPN de couche 3 ou BGP-MPLS VPN [rfc 2547]
 - les VPN de couche 2 [draft Martini et Kompella]

VPN de couche 3

- Définitions (RFC 2547) :
 - Deux sites distincts qui appartiennent à un même VPN sur un backbone
 - ont une connectivité IP sur ce même backbone
 - Intranet : les sites dans un VPN appartiennent à la même entreprise
 - Extranet : tous les sites dans un VPN appartiennent différentes entreprises
 - Un site peut être dans plusieurs VPN
 - Dans un intranet et dans un, ou plusieurs extranets
- Requièrent une grande coordination entre le client et le "provider" que l'approche VPN couche 2
- Ils sont actuellement les plus répandus au niveau des offres commerciales

Fonctionnement (1)

- Un identifiant RD (*Route Distinguisher*) est associé à chaque route distincte
 - permet d'affecter une adresse VPN-IPv4 ou VPN-IPv6 unique à chaque IP VPN
- Un LER à l'entrée va rajouter ce RD alors qu'un LER en sortie va le supprimer



- LER contient table VRF (*VPN Routing and Forwarding*)
 - Consultée par les paquets émis par un site faisant parti du VPN (du site émetteur)
 - Contient les routes reçues du CE ainsi que des autres LERs

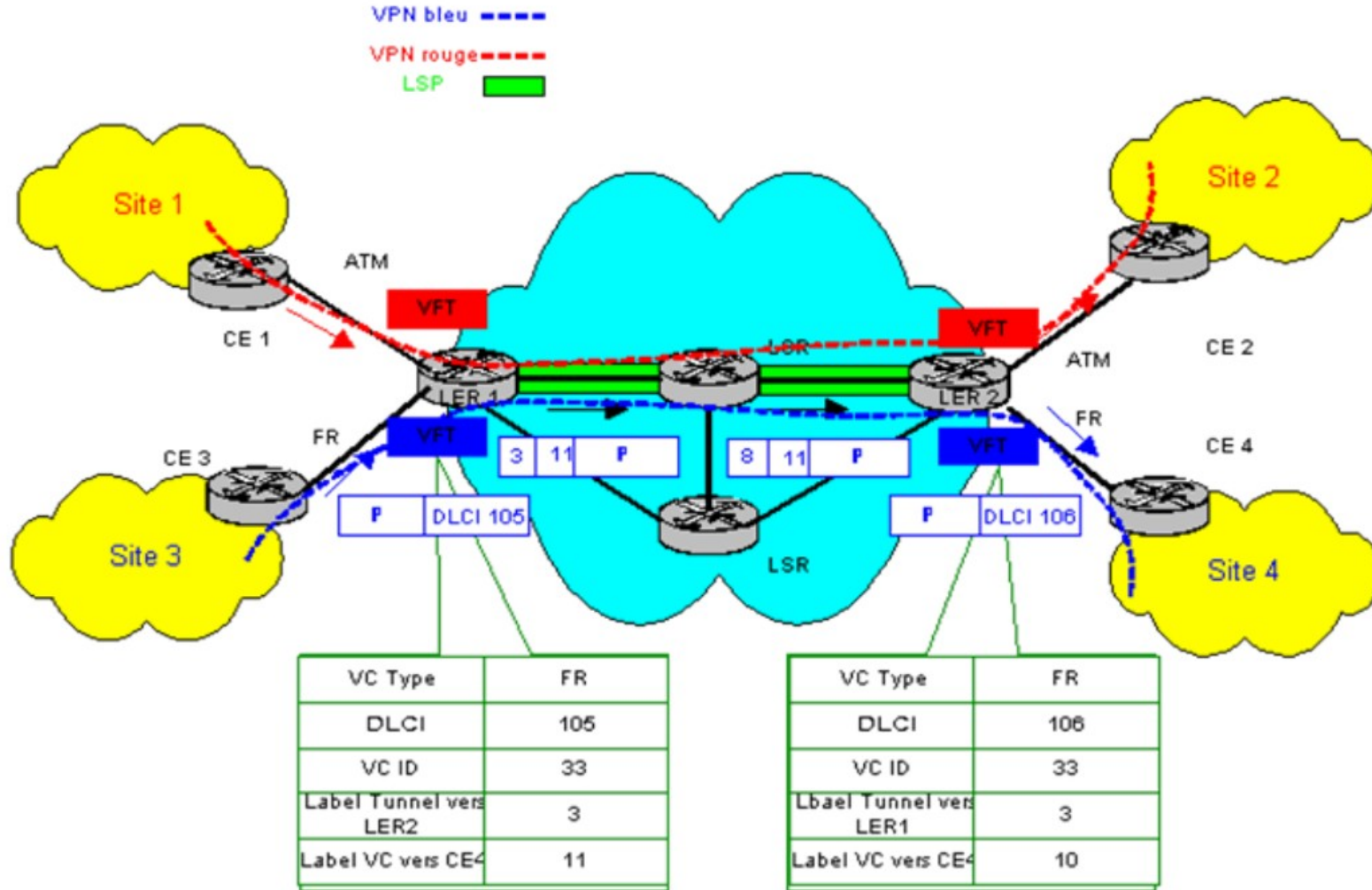
Fonctionnement (2)

- Information de routage entre CE et LER
 - distribuée via des protocoles de routages classiques
 - OSPF, RIP, BGP, routes statiques,...
- La création d'un tunnel VPN en MPLS exige de gérer une pile de profondeur 2 au niveau du paquet
 - 1er niveau de label
 - sert à gérer le prochain saut en terme de LER (routeur LER de sortie)
 - déterminé par un protocole style BGP
 - 2ème niveau de label
 - sert à gérer le prochain saut en terme de LSR dans le tunnel au sein du réseau MPLS
- Séparation des flux entre différents VPNs sur un même routeur MPLS
 - assurée par le fait que les tables VRF sont propres à chaque VPN
 - ne peuvent donc pas être consultées par les paquets associés à d'autres VPNs.

VPN de couche 2

- *Transparent LAN Services (TLS) ou Virtual Private LAN Services (VPLS)*
 - restent encore propriétaires mais tendent à se standardiser
 - l' IETF : *Martini drafts* et *Kompella drafts*
 - S'adressent aux fournisseurs d'accès qui ont peu de VPN à établir et qui veulent une indépendance au niveau du coeur du réseau
- S'établissent directement au niveau 2 dans le cas de réseau ATM ou Frame Relay
- Fournissent des solutions simples
- Permettent le transport de n'importe quel protocole de la couche 3
- Les LERs en entrée n'ont pas à gérer les tables de routages
 - sont moins chargés que dans le cas des VPN MPLS de niveau 3

Fonctionnement



■ LER

- ❑ fournit aux clients des *circuit IDs* de niveau 2 (VPI/VCI, DLCI, ...)
- ❑ associe ces *circuit IDs* a des LSP MPLS qui traversent le réseau

Le protocole OSPF et IS-IS

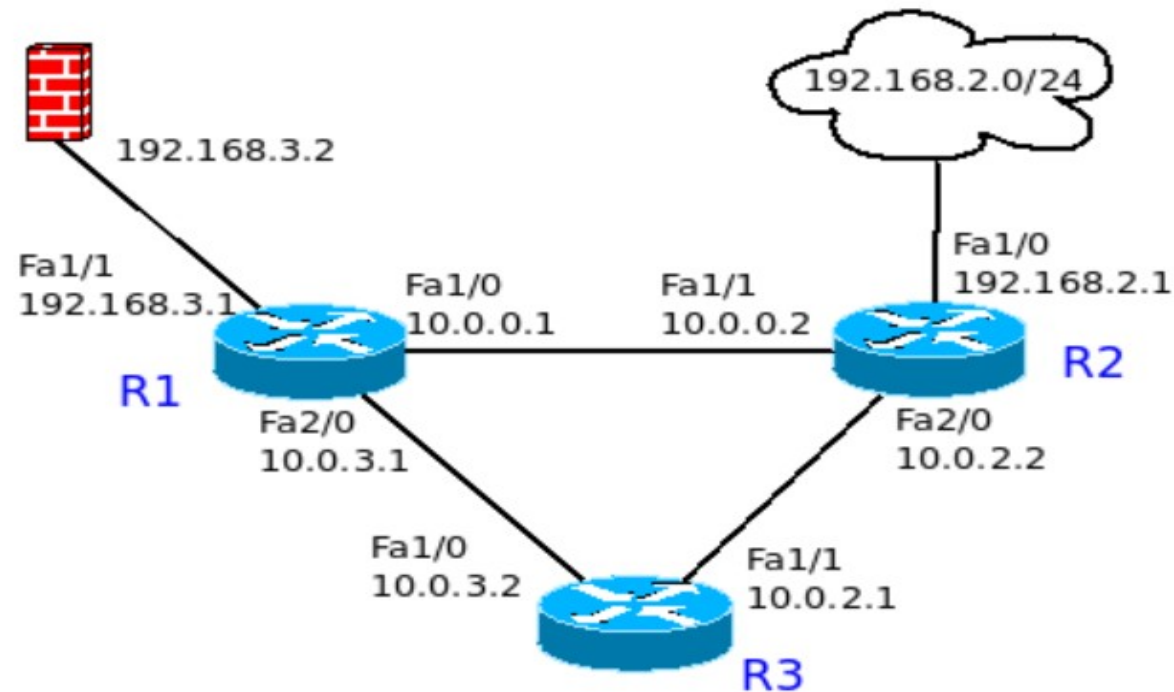
Le protocole OSPF

Définition

- OSPF est un protocole de routage à état de liens
- Les protocoles de routage à état de lien fonctionnent en envoyant **des publicités à état de lien (LSA)** à tous les autres routeurs à état de lien.
- Tous les routeurs doivent disposer de ces annonces d'état de lien afin de pouvoir créer leur **base de données d'état de lien** ou **LSDB** . Fondamentalement, toutes les publicités d'état de lien sont une pièce du puzzle qui construit la LSDB.
- Si vous disposez de nombreux routeurs OSPF, il n'est peut-être pas très efficace que chaque routeur OSPF transmette ses LSA à tous les autres routeurs OSPF. Permettez-moi de vous montrer un exemple:

Configuration OSPF

Exemple de configuration OSPF



Configuration OSPF

Comment configurer le protocole OSPF sur un routeur Cisco

Configuration IP de l'interface de loopback

```
R1(config)#interface loopback 0
```

```
R1(config-if)#ip address 172.16.0.1 255.255.255.0
```

Activation du protocole de routage OSPF et des réseaux participant aux annonces

Le numéro du processus ospf est 100.

L'adresse IP de l'interface connectée à un autre routeur est 10.0.3.1/24 et le numéro de l'aire est 0.

```
R1(config)#router ospf 100
```

```
R1(config-router)#network 10.0.3.0 0.0.0.255 area 0
```

```
R1(config-router)#network 10.0.0.0 0.0.0.255 area 0
```


Affichage des informations concernant ospf

R1#show ip ospf

Routing Process "ospf 100" with ID 172.16.0.1

Start time: 00:18:40.612, Time elapsed: 00:08:28.352

Supports only single TOS(TOS0) routes

Supports opaque LSA

Supports Link-local Signaling (LLS)

Supports area transit capability

Router is not originating router-LSAs with maximum metric

Initial SPF schedule delay 5000 msec

Minimum hold time between two consecutive SPF's 10000 msec

Maximum wait time between two consecutive SPF's 10000 msec

Afficher les voisins ospf

R1#show ip ospf neighbor

Neighbor ID	Pri	State	Dead Time	Address	Interface
172.16.0.3	1	FULL/BDR	00:00:39	10.0.3.2	FastEthernet2/0
172.16.0.2	1	FULL/DR	00:00:39	10.0.0.2	FastEthernet1/0

R1#show ip ospf neighbor detail

Neighbor 172.16.0.3, interface address 10.0.3.2
In the area 0 via interface FastEthernet2/0
Neighbor priority is 1, State is FULL, 34 state changes
DR is 10.0.3.1 BDR is 10.0.3.2
Options is 0x12 in Hello (E-bit, L-bit)
Options is 0x52 in DBD (E-bit, L-bit, O-bit)
LLS Options is 0x1 (LR)
Dead timer due in 00:00:35
Neighbor is up for 00:11:57

Redistribution de route dans ospf

Redistribution des routes connectées:

```
R2(config)#router ospf 100
```

```
R2(config-router)#redistribute connected
```

Redistribution d'une route par défaut:

Configuration d'une route par défaut puis redistribution dans ospf:

```
R4(config)#ip route 0.0.0.0 0.0.0.0 192.168.3.2
```

```
R1(config)#router ospf 100
```

```
R1(config-router)#redistribute static
```

```
R1(config-router)#default-information originate
```

Authentication

```
R1(config)#router ospf 100
```

```
R1(config-router)#area 0 authentication message-digest
```

```
R1(config-router)#exit
```

```
R1(config)#int fa 1/0
```

```
R1(config-if)#ip ospf message-digest-key 1 md5 password
```

```
R1(config-if)#int fa 2/0
```

```
R1(config-if)#ip ospf message-digest-key 1 md5 password
```

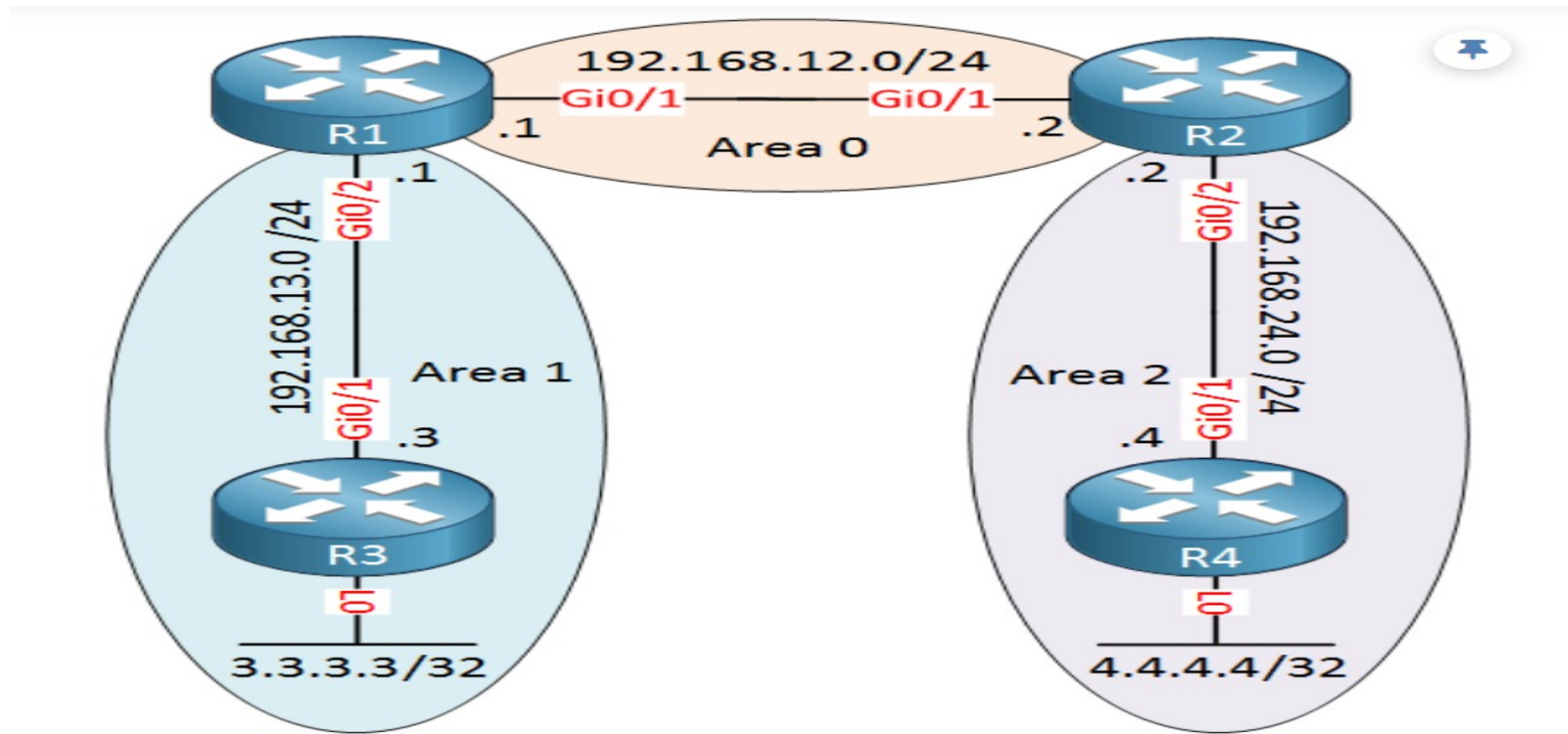
```
R1(config-if)#
```

Vérification

sh run sur chaque routeur pour vérifier la configuration

Configuration multizone OSPF

Nous utiliserons la topologie suivante :



Nous avons R1 et R2 dans la zone 0, la zone de base. Entre R1 et R3, nous utiliserons la zone 1 et entre R2/R4 nous utiliserons la zone 2. R3 et R4 ont une interface de bouclage avec une adresse IP que nous publierons dans leur zone.

Configuration

Commençons par toutes les commandes réseau pour que OSPF soit opérationnel. La commande network définit à quelle zone appartiendra chaque interface. Tout d'abord, nous allons configurer R1 et R2 pour la zone backbone :

Routeur 1:

```
R1(config)#router ospf 1
```

```
R1(config-router)#network 192.168.12.0 0.0.0.255 area 0
```

Routeur 2:

```
R2(config)#router ospf 1
```

```
R2(config-router)#network 192.168.12.0 0.0.0.255 area 0
```

Configurons R1 et R3 pour la zone 1 :

R1(config)#router ospf 1

R1(config-router)#network 192.168.13.0 0.0.0.255 area 1

R3(config)#router ospf 1

R3(config-router)#network 192.168.13.0 0.0.0.255 area 1

R3(config-router)#network 3.3.3.3 0.0.0.0 area 1

R2 et R4 pour la zone 2 :

R2(config)#router ospf 1

R2(config-router)#network 192.168.24.0 0.0.0.255 area 2

R4(config)#router ospf 1

R4(config-router)#network 192.168.24.0 0.0.0.255 area 2

R4(config-router)#network 4.4.4.4 0.0.0.0 area 2

Vérification

```
R1#show ip ospf neighbor
```

Neighbor ID	Pri	State	Dead Time	Address	Interface
192.168.24.2	1	FULL /DR	00:00:36	192.168.12.2	GigabitEthernet0/1
3.3.3.3	1	FULL /BDR	00:00:34	192.168.13.3	GigabitEthernet0/2

R1 a formé une contiguïté voisine avec R2 et R3. Vérifions R2 :

```
R2#show ip ospf neighbor
```

Neighbor ID	Pri	State	Dead Time	Address	Interface
192.168.13.1	1	FULL /BDR	00:00:34	192.168.13.1	GigabitEthernet0/1
4.4.4.4	1	FULL /BDR	00:00:30	192.168.24.4	GigabitEthernet0/2

R2 a formé des contiguïtés voisines avec R1 et R4. Cependant, la commande `show ip ospf voisin` ne me dit rien sur les zones utilisées. Si vous voulez voir cela, vous pouvez ajouter le paramètre **de détail** comme ceci :

Vérification

```
R2#show ip ospf neighbor detail
```

```
Neighbor 192.168.13.1, interface address 192.168.12.1
```

```
In the area 0 via interface GigabitEthernet0/1
```

```
Neighbor priority is 1, State is FULL, 6 state changes
```

```
DR is 192.168.12.2 BDR is 192.168.12.1
```

```
Options is 0x12 in Hello (E-bit, L-bit)
```

```
Options is 0x52 in DBD (E-bit, L-bit, O-bit)
```

```
LLS Options is 0x1 (LR)
```

```
Dead timer due in 00:00:33
```

```
Neighbor is up for 00:17:30
```

```
Index 1/1/1, retransmission queue length 0, number of retransmission 0
```

```
First 0x0(0)/0x0(0)/0x0(0) Next 0x0(0)/0x0(0)/0x0(0)
```

```
Last retransmission scan length is 0, maximum is 0
```

```
Last retransmission scan time is 0 msec, maximum is 0 msec
```

```
Neighbor 4.4.4.4, interface address 192.168.24.4
```

```
In the area 2 via interface GigabitEthernet0/2
```

```
Neighbor priority is 1, State is FULL, 6 state changes
```

```
DR is 192.168.24.2 BDR is 192.168.24.4
```

```
Options is 0x12 in Hello (E-bit, L-bit)
```

```
Options is 0x52 in DBD (E-bit, L-bit, O-bit)
```

```
LLS Options is 0x1 (LR)
```

```
Dead timer due in 00:00:31
```

```
Neighbor is up for 00:15:57
```

Vérification

```
Index 1/1/2, retransmission queue length 0, number of retransmission 0  
First 0x0(0)/0x0(0)/0x0(0) Next 0x0(0)/0x0(0)/0x0(0)  
Last retransmission scan length is 0, maximum is 0  
Last retransmission scan time is 0 msec, maximum is 0 msec
```

NB: vous pouvez voir que l'interface GigabitEthernet0/1 est dans la zone 0 et l'interface GigabitEthernet0/2 est dans la zone 2.

Verification

Une autre bonne commande pour trouver des informations sur la zone est **show ip protocols** :

```
R2#show ip protocols
*** IP Routing is NSF aware ***

Routing Protocol is "application"
  Sending updates every 0 seconds
  Invalid after 0 seconds, hold down 0, flushed after 0
  Outgoing update filter list for all interfaces is not set
  Incoming update filter list for all interfaces is not set
  Maximum path: 32
  Routing for Networks:
  Routing Information Sources:
    Gateway          Distance      Last Update
  Distance: (default is 4)

Routing Protocol is "ospf 1"
  Outgoing update filter list for all interfaces is not set
  Incoming update filter list for all interfaces is not set
  Router ID 192.168.24.2
  It is an area border router
  Number of areas in this router is 2. 2 normal 0 stub 0 nssa
  Maximum path: 4
```

Routing for Networks:

192.168.12.0 0.0.0.255 area 0

192.168.24.0 0.0.0.255 area 2

Routing Information Sources:

Gateway	Distance	Last Update
4.4.4.4	110	00:16:04
192.168.13.1	110	00:16:53

Distance: (default is 110)

NB: vous pouvez voir quels réseaux appartiennent à quelle zone :

- Réseau 192.168.12.0 en zone 0.
- Réseau 192.168.24.0 en zone 2.

Vérification: tables de routage

```
R1#show ip route ospf
```

```
Codes: L - local, C - connected, S - static, R - RIP, M - mobile, B - BGP
```

```
D - EIGRP, EX - EIGRP external, O - OSPF, IA - OSPF inter area
```

```
N1 - OSPF NSSA external type 1, N2 - OSPF NSSA external type 2
```

```
E1 - OSPF external type 1, E2 - OSPF external type 2
```

```
i - IS-IS, su - IS-IS summary, L1 - IS-IS level-1, L2 - IS-IS level-2
```

```
ia - IS-IS inter area, * - candidate default, U - per-user static route
```

```
o - ODR, P - periodic downloaded static route, H - NHRP, l - LISP
```

```
a - application route
```

```
+ - replicated route, % - next hop override, p - overrides from PfR
```

```
Gateway of last resort is not set
```

```
3.0.0.0/32 is subnetted, 1 subnets
```

```
O      3.3.3.3 [110/2] via 192.168.13.3, 00:01:47, GigabitEthernet0/2
```

```
4.0.0.0/32 is subnetted, 1 subnets
```

```
O IA    4.4.4.4 [110/3] via 192.168.12.2, 00:00:54, GigabitEthernet0/1
```

```
O IA 192.168.24.0/24 [110/2] via 192.168.12.2, 00:01:44, GigabitEthernet0/1
```

NB: Nous voyons trois entrées OSPF. Le premier concerne la 3.3.3.3/32, l'interface de bouclage de R3. Il apparaît avec un O puisqu'il s'agit d'une route intra-zone. R1 a également pris connaissance des versions 4.4.4.4/32 et 192.168.24.0/24. Ces deux entrées apparaissent comme O IA puisqu'il s'agit de routes inter-zones.

```
R2#show ip route ospf
```

```
      3.0.0.0/32 is subnetted, 1 subnets
O IA    3.3.3.3 [110/3] via 192.168.12.1, 00:02:19, GigabitEthernet0/1
      4.0.0.0/32 is subnetted, 1 subnets
O        4.4.4.4 [110/2] via 192.168.24.4, 00:01:29, GigabitEthernet0/2
O IA    192.168.13.0/24 [110/2] via 192.168.12.1, 00:02:24, GigabitEthernet0/1
```

Nous voyons que R2 a appris 3.3.3.3/32 et 192.168.13.0/24 quelles sont les routes inter-zones. 4.4.4.4/32 est une route intra-zone.

Vérification

```
R3#show ip route ospf
```

```
4.0.0.0/32 is subnetted, 1 subnets
```

```
O IA    4.4.4.4 [110/4] via 192.168.13.1, 00:01:57, GigabitEthernet0/1
O IA    192.168.12.0/24 [110/2] via 192.168.13.1, 00:02:50, GigabitEthernet0/1
O IA    192.168.24.0/24 [110/3] via 192.168.13.1, 00:02:47, GigabitEthernet0/1
```

Tout ce que R3 a appris vient d'une autre zone, c'est pourquoi nous ne voyons ici que les itinéraires inter-zones. La même chose s'applique à R4 :

```
R4#show ip route ospf
```

```
3.0.0.0/32 is subnetted, 1 subnets
```

```
O IA    3.3.3.3 [110/4] via 192.168.24.2, 00:02:13, GigabitEthernet0/1
O IA    192.168.12.0/24 [110/2] via 192.168.24.2, 00:02:13, GigabitEthernet0/1
O IA    192.168.13.0/24 [110/3] via 192.168.24.2, 00:02:13, GigabitEthernet0/1
```

Vérification

Juste pour être sûr, essayons un ping rapide entre R3 et R4 pour prouver que notre configuration OSPF multizone fonctionne :

```
R3#ping 4.4.4.4 source 3.3.3.3
```

```
Type escape sequence to abort.
```

```
Sending 5, 100-byte ICMP Echos to 4.4.4.4, timeout is 2 seconds:
```

```
Packet sent with a source address of 3.3.3.3
```

```
!!!!
```

```
Success rate is 100 percent (5/5), round-trip min/avg/max = 9/11/13 ms
```

Le protocole IS-IS:

Fonctionnement de IS-IS

- Intermediate **S**ystem to Intermediate **S**ystem
- ISO 10589 spécifie le protocole de routage IS-IS dans le modèle OSI pour le trafic CLNS
 - Est un protocole à état de lien avec deux niveaux d'architecture hiérarchique
- RFC 1195 a ajouté le support de IP
 - Integrated IS-IS
 - I/IS-IS fonctionne au dessus de la couche Data Link

Fonctionnement de IS-IS

- Les routeurs avec IS-IS activé recherchent les routeurs voisins qui exécutent également IS-IS
 - Hello Protocol Data Units (PDUs) sont échangés
 - The “Hello” packet inclus la liste des voisins connus et les details tels que “Interval Hello” et “Router dead interval”
 - Hello interval – A quelle fréquence le routeur enverra les paquets Hello
 - Router dead interval – Quel est le temps d’attente avant de decider que le routeur n’est plus disponible
 - Les valeurs de “hello interval” et “router dead interval” doivent concorder sur les voisins.
 - Lorsqu’un voisin repond avec des valeurs concordant aux valeurs du packet Hello, l’adjacence se forme.

Relations de Voisinage IS-IS

- Une relation est établie entre des routeurs voisins dans le but d'échanger des informations de routage: Cela est appelée Adjacence/Voisinage.

Adjacences IS-IS

- Une fois l'adjacence formée, les voisins partagent les informations sur l'état des liens.
 - L'information est envoyée à travers le Link State PDU (LSP)
 - Les LSPs sont envoyées à tous les voisins
- Une nouvelle information reçue d'un voisin est utilisée pour créer une nouvelle vue de la topologie du Réseau.
- Si une liaison est indisponible
 - De nouveaux LSPs sont envoyées
 - Les routeurs reconstruisent la table de routage

Niveaux IS-IS

- IS-IS a deux couches hiérarchiques
 - Level-2 (Pour le BackBone)
 - Level-1 (Pour l'accès)
- Un routeur peut être
 - Level-1 (L1) router
 - Level-2 (L2) router
 - Level-1-2 (L1/L2) router

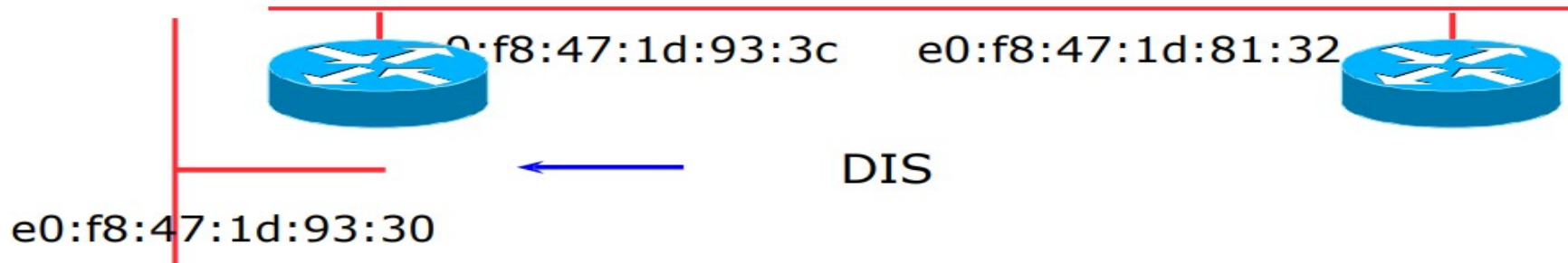
Liens dans IS-IS

- Deux types de liens IS-IS:
 - Lien Point-to-Point
 - Seulement deux routeurs sur le lien pour former l'adjacence.
 - Multi-access network (c-à-d. ethernet)
 - Plusieurs routeurs dans le Réseau avec plusieurs adjacences.
- IS-IS dans les réseaux multi-access networks nécessite des optimisations pour faciliter la convergence
 - Un routeur est élu pour l'envoi des LSPs à tous les autres routeurs. Il est appelé "**Designated Information System**"
 - Les routeurs dans le Réseau Multi-access forment l'adjacence avec le routeur DIS

Sélection du Routeur Désigné

- Configuration de la priorité (par interface)
 - Configurer la plus haute priorité sur le routeur pour qu'il soit le DIS

```
interface gigabitethernet0/1
isis priority 127 level-2
```
- Sinon la priorité est déterminée par la plus grande MAC Address
 - Les bonnes pratiques exigent de mettre au moins deux routeurs souhaités avec les grandes valeurs pour avoir un DIS primaire et secondaire.



Adjacences: Exemples

- Pour voir l'état d'une adjacence CLNS:

```
show clns neighbor
```

System Id	Interface	SNPA	State	Holdtime	Type	Protocol
Router2	Fa0/0	ca01.9798.0008	Up	23	L2	M-ISIS
Router3	Se1/0	*HDLC*	Up	26	L2	M-ISIS

- Pour voir l'état de l'adjacence IS-IS:

```
show isis neighbor
```

System Id	Type	Interface	IP Address	State	Holdtime	Circuit Id
Router2	L2	Fa0/0	10.10.15.2	UP	24	Router2.01
Router3	L2	Se1/0	10.10.15.6	UP	27	00

IS-IS sur Cisco IOS

- Demarrer les configurations IS-IS sur CISCO par la commande

```
router isis as42
```

- ou “as42” est l’ID de processus (Process ID)
- IS-IS process ID est unique sur le routeur
 - Donne la possibilité de lancer plusieurs instances de IS-IS
 - Le process ID n’est pas utilise pour la communication entre routeurs.
 - Quelques ISPs utilise la valeur de leur AS BGP comme process ID de IS-IS

IS-IS sur Cisco IOS

- Les routeurs Cisco ont par défaut le niveau L1/L2
- Une fois IS-IS démarré, d'autres configurations sont requises sous le process IS-IS:
 - Afficher les changements d'adjacence dans les logs
`systemes`
`log-adjacency-changes`
 - Configure le **metric-style** à la valeur **wide**
`metric-style wide`
 - Définir le type d'IS au niveau 2
`is-type level-2-only`
 - Configurer le NET address
`net`
`49.0001.<loopback>.00`

Adresse NSAP IS-IS

- Les protocoles de routage IP disposent d'un **router-id** pour identifier de façon unique un routeur.
- IS-IS utilise l'adresse NSAP
 - Peut être de taille variable entre 64 et 160 bits
- ISPs typically choose NSAP addresses thus:
 - 8 premiers bits – Choisissez un nombre (généralement la valeur 49)
 - Les prochains 16 bits définissent l'**area**
 - Les prochains 48 bits définissent le **system ID**
 - Les derniers 8 bits sont définis à zéro
- Example:
 - NSAP: 49.0001.1921.6800.1001.00
 - Router: 192.168.1.1 (loopback) in Area 1

Ajout d'interfaces dans IS-IS

- Pour activer l'IS-IS sur une interface:

```
interface POS4/0  
  ip router isis as42
```

- Met l'adresse de sous-Réseau dans la LSDB
- Active le CLNS sur cette interface

- Pour désactiver IS-IS sur une interface:

```
router isis as42  
  passive-interface GigabitEthernet 0/0
```

- Désactive CLNS sur une interface

Authentification de voisins IS-IS

- Authentification entre voisins est hautement recommandée
 - Empêche les routeurs non autorisés pour la formation d'adjacence.

- Création d'une **key-chain**

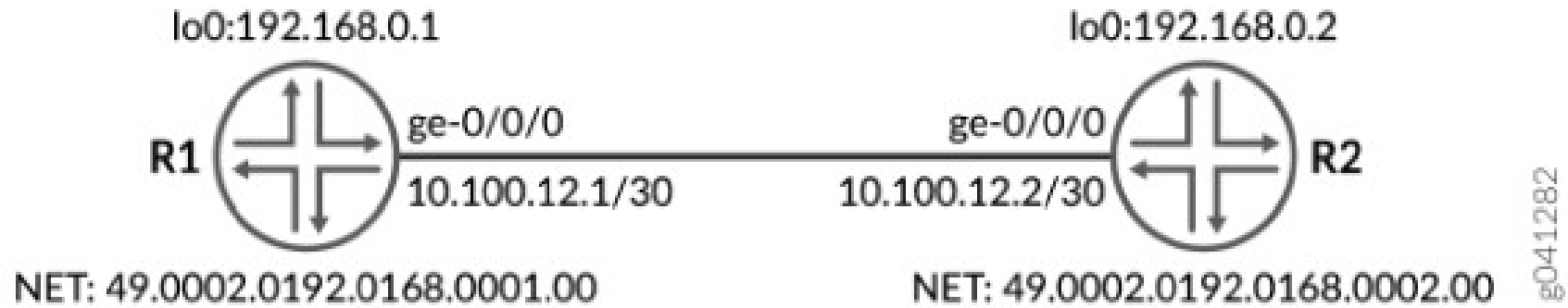
```
key chain isis-as42
  key 1
    key-string <password>
```

!

- Appliquer une key-chain à une interface

```
interface POS 4/0
  isis authentication mode md5 level-2
  isis authentication key-chain isis-as42 level-2
```

Architecture



Implémentation de IS-IS

1) Configuration rapide

- R1:

```
set security forwarding-options family iso mode packet-based
set interfaces ge-0/0/0 unit 0 description "To R2"
set interfaces ge-0/0/0 unit 0 family inet address 10.100.12.1/30
set interfaces ge-0/0/0 unit 0 family iso
set interfaces lo0 unit 0 family inet address 192.168.0.1/32
set interfaces lo0 unit 0 family iso address 49.0002.0192.0168.0001.00
set protocols isis interface ge-0/0/0.0
set protocols isis interface lo0.0
```


1) Configuration rapide

- R2:

```
set security forwarding-options family iso mode packet-based
set interfaces ge-0/0/0 unit 0 description "To R1"
set interfaces ge-0/0/0 unit 0 family inet address 10.100.12.2/30
set interfaces ge-0/0/0 unit 0 family iso
set interfaces lo0 unit 0 family inet address 192.168.0.2/32
set interfaces lo0 unit 0 family iso address 49.0002.0192.0168.0002.00
set protocols isis interface ge-0/0/0.0
set protocols isis interface lo0.0
```

Procédure étape par étape

- Pour configurer IS-IS :

1. (Sur les équipements de sécurité uniquement) Activez le transfert du trafic IS-IS pour surmonter le comportement par défaut de suppression du trafic IS-IS.

```
[edit security]
user@R1# set forwarding-options family iso mode packet-based
```

2. Configurez l'interface de l'équipement et activez la famille ISO sur l'interface.

```
[edit interfaces]
user@R1# set ge-0/0/0 unit 0 description "To R2"
user@R1# set ge-0/0/0 unit 0 family inet address 10.100.12.1/30
user@R1# set ge-0/0/0 unit 0 family iso
```

Procédure étape par étape: suite

3. Configurez l'interface de bouclage et définissez l'adresse NET.

```
[edit interfaces]
user@R1# set lo0 unit 0 family inet address 192.168.0.1/32
user@R1# set lo0 unit 0 family iso address 49.0002.0192.0168.0001.00
```

4. Activez IS-IS pour les interfaces des équipements.

```
[edit protocols]
user@R1# set isis interface ge-0/0/0.0
user@R1# set isis interface lo0.0
```

NB: Si vous avez fini de configurer l'équipement, saisissez `commit` à partir du mode de configuration.

Vérification

```
user@R1> show isis adjacency
```

Interface	System	L State	Hold (secs)	SNPA
ge-0/0/0.0	R2	1 Up	7	56:4:1e:0:5f:58
ge-0/0/0.0	R2	2 Up	7	56:4:1e:0:5f:58

```
user@R1> show isis interface
```

IS-IS interface database:

Interface	L	CirID	Level 1	DR	Level 2	DR	L1/L2	Metric
ge-0/0/0.0	3	0x1	R2.02		R2.02		10/10	
lo0.0	3	0x1	Passive		Passive		0/0	

```
user@R1> ping 192.168.0.1 source 192.168.0.2 count 100 rapid
```

```
PING 192.168.0.2 (192.168.0.2): 56 data bytes
```

.....

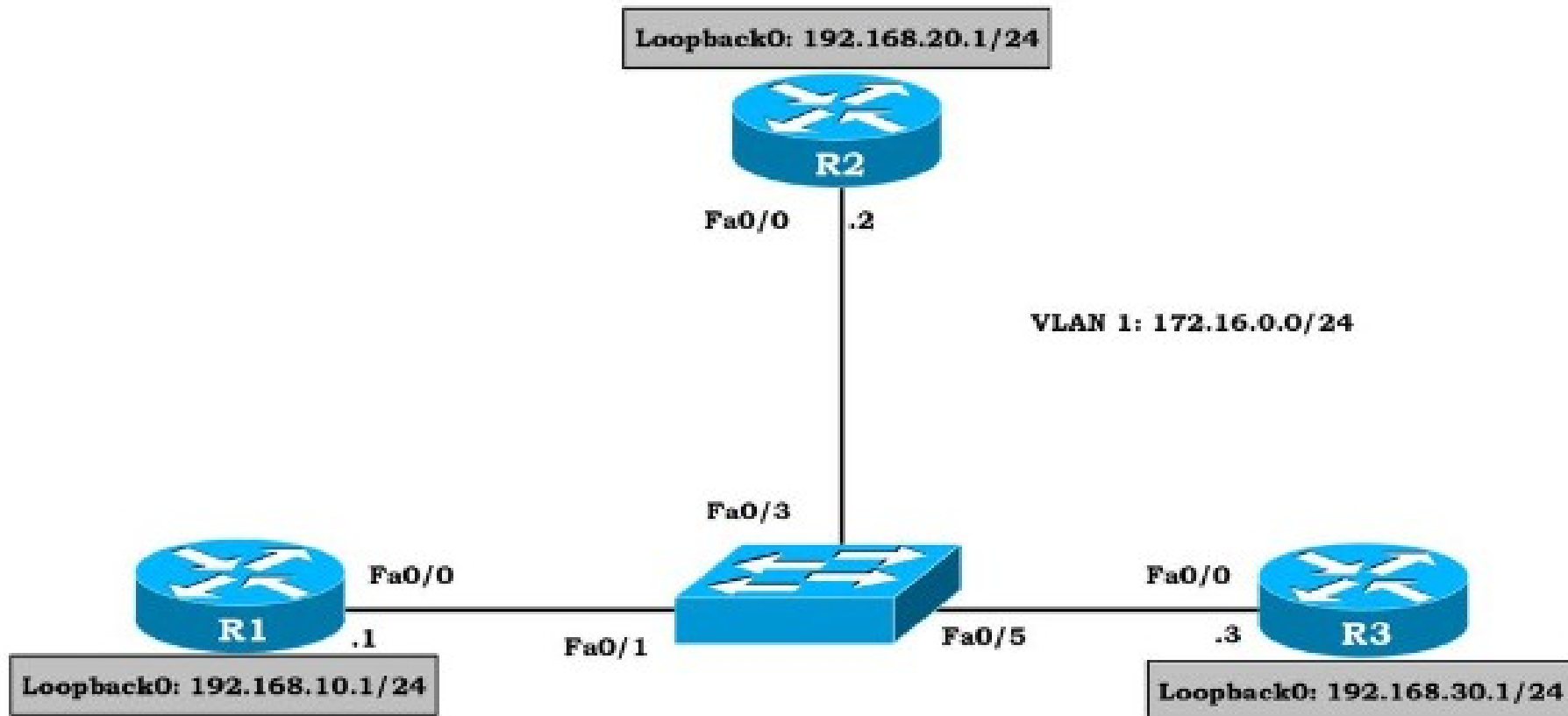
```
--- 192.168.0.2 ping statistics ---
```

```
100 packets transmitted, 100 packets received, 0% packet loss
```

```
round-trip min/avg/max/stddev = 1.807/2.425/6.447/0.553 ms
```

Configuration de Base IS-IS

Schéma du réseau



IS-IS (ISO/IEC 10589) est implémenté avec des adresses de point de service réseau (NSAP) constituées de trois champs: adresse de zone (Area), Système ID et NSEL (connu également comme N-selector, identifieur du service ou ID du processus). Le champ adresse de zone peut avoir de un à treize octets, le système ID a usuellement six octets (doit avoir six octets pour l'IOS Cisco) et NSEL identifie un processus sur un équipement. C'est un peu une équivalence avec le port ou le socket dans IP. Le champ NSEL n'est pas utilisé dans les décisions de routage.

Quand le champ NSEL est fixé à 00, le NSAP est référencé comme Network Entity Title ou NET. Les NETs et les NSAPs sont représentés en hexadécimal et doivent commencer et se terminer à une limite d'octet tel 49.0001.1111.1111.1111.00.

Le routage IS-IS Level 1 ou L1 est basé sur le "System ID". par conséquent chaque routeur doit avoir un System ID unique dans la zone. Le routage IS-IS L1 équivaut à un routage intra-area. Il est de coutume d'utiliser soit une adresse MAC du routeur ou pour IS-IS intégré de coder l'adresse IP d'un interface loopback dans le system ID.

Les adresses de zone commençant par 48, 49, 50 ou 51 sont des adresses privées. Ce groupe d'adresses ne doit pas être annoncé aux autres réseaux CLNS (ConnectionLess Network Service). L'adresse de zone (area) doit être la même pour tous les routeurs dans l'area.

Sur un LAN, un des routeurs est élu DIS (Designated Intermediate System) sur la base de la priorité d'interface. La valeur par défaut est 64. Si toutes les interfaces ont la même priorité, le routeur avec l'adresse de SNPA (SubNetwork Point of Attachment) la plus élevée est choisi. L'adresse MAC sert d'adresse SNPA pour les réseaux LAN Ethernet. Le DIS a le même rôle que celui du routeur désigné dans OSPF. L'ingénieur réseau de l'ITA décide que le routeur R1 est le DIS aussi sa priorité devra être plus élevée que celles de R2 et de R3.

Routeur 1

Configurer le protocole IS-IS sur chaque routeur et fixez la priorité à 100 pour l'interface FastEthernet 0/0 de R1

- R1 (config)# routeur isis
- R1 (config-routeur)# net 49.0001.1111.1111.1111.00
- R1 (config-routeur)# interface fastethernet 0/0
- R1 (config-if)# ip routeur isis
- R1 (config-if)# isis priority 100
- R1 (config-if)# interface loopback 0
- R1 (config-if)# ip routeur isis

Routeur 2

- R2 (config)# routeur isis
- R2 (config-routeur)# net 49.0001.2222.2222.2222.00
- R2 (config-routeur)# interface fastethernet 0/0
- R2 (config-if)# ip routeur isis
- R2 (config-if)# interface loopback 0
- R2 (config-if)# ip routeur isis

Routeur 3

- R3 (config)# routeur isis
- R3 (config-routeur)# net 49.0001.3333.3333.3333.00
- R3 (config-routeur)# interface fastethernet 0/0
- R3 (config-if)# ip routeur isis
- R3 (config-if)# interface loopback 0
- R3 (config-if)# ip routeur isis

Identification des paramètres

- Adresse d'area
 - System ID R1
 - System ID R2
 - System ID R3
 - NSEL
-
- Pour cela, nous allons effectuer la commande suivante sur chaque routeur

Identification des paramètres

- Show ip protocols
- Show cnls protocols
- Show cnls neighbors

```
R1# show ip protocols
Routing Protocol is "isis"
  Invalid after 0 seconds, hold down 0, flushed after 0
```

```
R1# show cnls protocols
IS-IS Router: <Null Tag>
  System Id: 1111.1111.1111.00 IS-Type: level-1-2
  Manual area address(es):
```

```
R1# show cnls neighbors
```

System Id	Interface	SNPA	State	Holdtime	Type	Protocol
R2	Fa0/0	0004.9ad2.d0c0	Up	9	L1L2	IS-IS
R3	Fa0/0	0002.16f4.1ba0	Up	29	L1L2	IS-IS

Identification des paramètres

- Show cnls neighbors details

```
R1# show cnls neighbors detail
```

System Id	Interface	SNPA	State	Holdtime	Type	Protocol
R2	Fa0/0	0004.9ad2.d0c0	Up	24	L1L2	IS-IS
Area Address(es): 49.0001						
IP Address(es): 172.16.0.2*						
Uptime: 00:07:30						
NSF capable						
R3	Fa0/0	0002.16f4.1ba0	Up	27	L1L2	IS-IS
Area Address(es): 49.0001						
IP Address(es): 172.16.0.3*						
Uptime: 00:07:00						
NSF capable						

- Show isis database

```
R1# show isis database
```

IS-IS Level-1 Link State Database:

LSPID		LSP Seq Num	LSP Checksum	LSP Holdtime	ATT/P/OL
R1.00-00	*	0x00000008	0x088F	1191	0/0/0
R1.01-00	*	0x00000002	0x9B60	1192	0/0/0
R2.00-00		0x00000001	0x8736	1190	0/0/0
R3.00-00		0x00000002	0x39A1	1195	0/0/0

IS-IS Level-2 Link State Database:

LSPID		LSP Seq Num	LSP Checksum	LSP Holdtime	ATT/P/OL
R1.00-00	*	0x00000017	0x4E1B	1195	0/0/0
R1.01-00	*	0x00000002	0x4D37	1192	0/0/0
R2.00-00		0x00000010	0xF4B9	1191	0/0/0
R3.00-00		0x00000002	0xD703	1195	0/0/0

Identification des paramètres

- Show cnls interface fastethernet 0/0

```
R1# show cnls interface fastethernet 0/0
FastEthernet0/0 is up, line protocol is up
Checksums enabled, MTU 1497, Encapsulation SAP
ERPDUs enabled, min. interval 10 msec.
CLNS fast switching enabled
CLNS SSE switching disabled
DEC compatibility mode OFF for this interface
Next ESH/ISH in 8 seconds
Routing Protocol: IS-IS
  Circuit Type: level-1-2
  Interface number 0x0, local circuit ID 0x1
  Level-1 Metric: 10, Priority: 100, Circuit ID: R1.01
  DR ID: R1.01
  Level-1 IPv6 Metric: 10
  Number of active level-1 adjacencies: 2

  Level-2 Metric: 10, Priority: 100, Circuit ID: R1.01
  DR ID: R1.01
  Level-2 IPv6 Metric: 10
  Number of active level-2 adjacencies: 2
  Next IS-IS LAN Level-1 Hello in 803 milliseconds
  Next IS-IS LAN Level-2 Hello in 2 seconds
```

Notez que le circuit ID est R1.01, lequel est composé du System ID et du pseudonode ID, identifie le DIS. Les informations Circuit Types, Levels, Metric et Priority sont également affichées.

Identification des paramètres

- Show isis database R1.00-00 detail

```
R1# show isis database R1.00-00 detail

IS-IS Level-1 LSP R1.00-00
LSPID          LSP Seq Num    LSP Checksum  LSP Holdtime  ATT/P/OL
R1.00-00      * 0x0000000B  0x0292        831           0/0/0
  Area Address: 49.0001
  NLPID: 0xCC
  Hostname: R1
  IP Address: 192.168.10.1
  Metric: 10 IP 172.16.0.0 255.255.255.0
  Metric: 10 IP 192.168.10.0 255.255.255.0
  Metric: 10 IS R1.02
  Metric: 10 IS R1.01

IS-IS Level-2 LSP R1.00-00
LSPID          LSP Seq Num    LSP Checksum  LSP Holdtime  ATT/P/OL
R1.00-00      * 0x0000000D  0x4703        709           0/0/0
  Area Address: 49.0001
  NLPID: 0xCC
  Hostname: R1
  IP Address: 192.168.10.1
  Metric: 10 IS R1.02
  Metric: 10 IS R1.01
  Metric: 20 IP 192.168.30.0 255.255.255.0
  Metric: 10 IP 192.168.10.0 255.255.255.0
  Metric: 10 IP 172.16.0.0 255.255.255.0
  Metric: 20 IP 192.168.20.0 255.255.255.0
```

La métrique IS-IS par défaut est égale à 10 pour chaque liaison mais notez que les métriques pour les réseaux 192.168.20.0 et 192.168.30.0 sont toutes les deux égales à 20. Ceci est dû au fait que ces réseaux ne sont pas directement connectés mais sont connectés sur les routeurs voisins.

Identification des paramètres

```
R1# show isis topology
```

IS-IS paths to level-1 routers

System Id	Metric	Next-Hop	Interface	SNPA
R1	--			
R2	10	R2	Fa0/0	0004.9ad2.d0c0
R3	10	R3	Fa0/0	0002.16f4.1ba0

IS-IS paths to level-2 routers

System Id	Metric	Next-Hop	Interface	SNPA
R1	--			
R2	10	R2	Fa0/0	0004.9ad2.d0c0
R3	10	R3	Fa0/0	0002.16f4.1ba0

Les entrées surlignées de la colonne SNPA sont les adresses MAC des interfaces FastEthernet0/0 des routeurs R2 et R3.

Entrez la commande **show isis route** pour voir la table de routage IS-IS L1:

```
R1# show isis route
```

```
IS-IS not running in OSI mode (*) (only calculating IP routes)
(*) Use "show isis topology" command to display paths to all routers
```

Cette commande n'est pas très utile car elle est spécifique au routage OSI. Rappelez-vous que IP IS-IS a été validé sur chaque routeur. Si CLNP avait été configuré dans le réseau, des informations plus intéressantes auraient été affichées.

Entrez la commande **show clns route** pour voir la table de routage IS-IS L2:

```
R1# show clns route
```

```
Codes: C - connected, S - static, d - DecnetIV
       I - ISO-IGRP, i - IS-IS, e - ES-IS
       B - BGP, b - eBGP-neighbor
```

```
C 49.0001.1111.1111.1111.00 [1/0], Local IS-IS NET
C 49.0001 [2/0], Local IS-IS Area
```

De nouveau cette commande n'est pas très utile car elle s'applique au routage OSI et non au routage IP.

Identification des paramètres

```
R1# show ip route  
<partie supprimée>
```

```
Gateway of last resort is not set
```

```
i L1 192.168.30.0/24 [115/20] via 172.16.0.3, FastEthernet0/0  
C    192.168.10.0/24 is directly connected, Loopback0  
    172.16.0.0/24 is subnetted, 1 subnets  
C      172.16.0.0 is directly connected, FastEthernet0/0  
i L1 192.168.20.0/24 [115/20] via 172.16.0.2, FastEthernet0/0
```

Le protocole BGP

Définition

Le protocole BGP (Border Gateway Protocol) est un protocole évolutif de routage dynamique qui est utilisé par des groupes de routeurs sur Internet pour partager des informations de routage.

L'objectif principal de BGP est de fournir un système de routage inter-domaines qui garantit l'échange sans boucle d'informations de routage entre systèmes autonomes.

Il existe une distinction entre un système autonome ordinaire et un système configuré avec BGP pour mettre en œuvre une politique de transit. Ce dernier est appelé FAI ou fournisseur de services.

BGP-4 présente de nombreuses améliorations par rapport aux protocoles antérieurs. Il est aujourd'hui largement utilisé sur Internet pour connecter les FAI et interconnecter les entreprises avec les FAI.

Utilisation du BGP4, évite que la table de routage Internet ne devienne trop volumineuse pour interconnecter des millions d'utilisateurs.

- ☐ Permet d'utiliser pleinement toute la bande passante en manipulant les attributs de chemin.
- ☐ Permet à un système autonome de contrôler le flux de trafic à l'aide de plusieurs attributs BGP
- ☐ Examinent le coût du chemin
- ☐ Choisissent le meilleur chemin d'un point d'un réseau d'entreprise à un autre en fonction de certaines mesures

Le protocole BGP

Border Gateway Protocol

- Protocole de routage utilisé pour échanger des informations de routage entre les différents réseaux
 - Exterior Gateway Protocol
- Décrit dans la RFC 4271
 - RFC4276 donne un rapport d'exécution sur BGP
 - RFC4277 décrit les expériences opérationnelles d'utilisation de BGP
- Le système autonome (AS) est la pierre angulaire de BGP
 - Il est utilisé pour identifier de façon unique les réseaux ayant une politique de routage commune (généralement 1 par entreprise)

Le protocole BGP: ROLES

BGP

- Path Vector Protocol
- Mises à jour incrémentales
- Beaucoup d'options pour l'application des stratégies
- Classless Inter Domain Routing (CIDR)
- Largement utilisé pour le backbone Internet
- Systèmes autonomes

Le protocole BGP

Path Vector Protocol

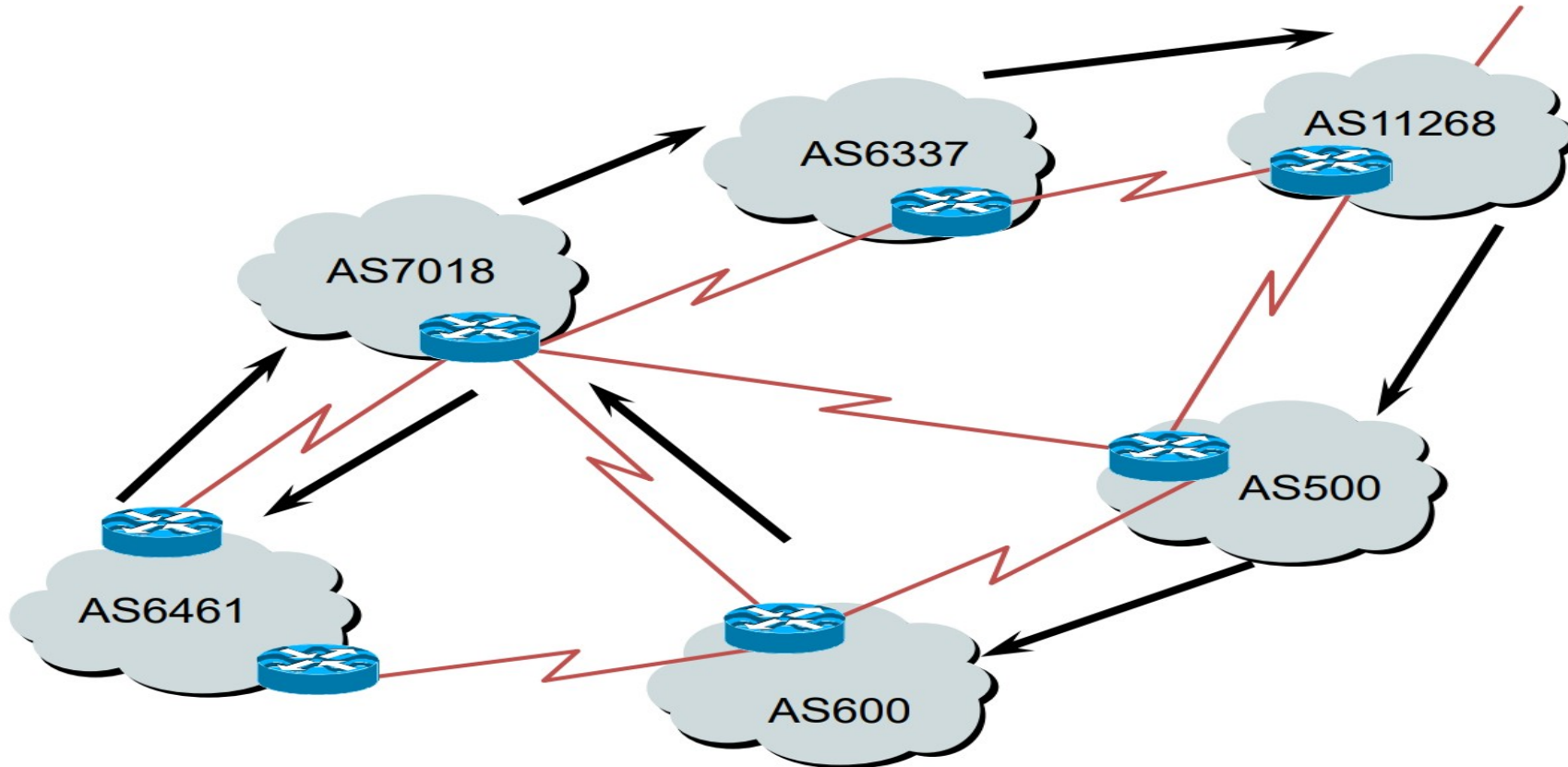
- BGP est classé comme un protocole de routage *path vector* (voir RFC 1322)
 - Un protocole path vector définit un trajet a partir d'une destination et les caractéristiques du PATH qui mène à cette destination.

12.6.126.0/24 207.126.96.43 1021 0 6461 7018 6337 11268 i

AS Path

Path Vector Protocol

Path Vector Protocol



Définitions

- **Transit** – acheminement de trafic sur un réseau, généralement moyennant un paiement
- **Peering** – échanger des informations de routage et de trafic gratuitement
- **Default** – où envoyer le trafic quand il n'y a pas de correspondance explicite dans la table de routage

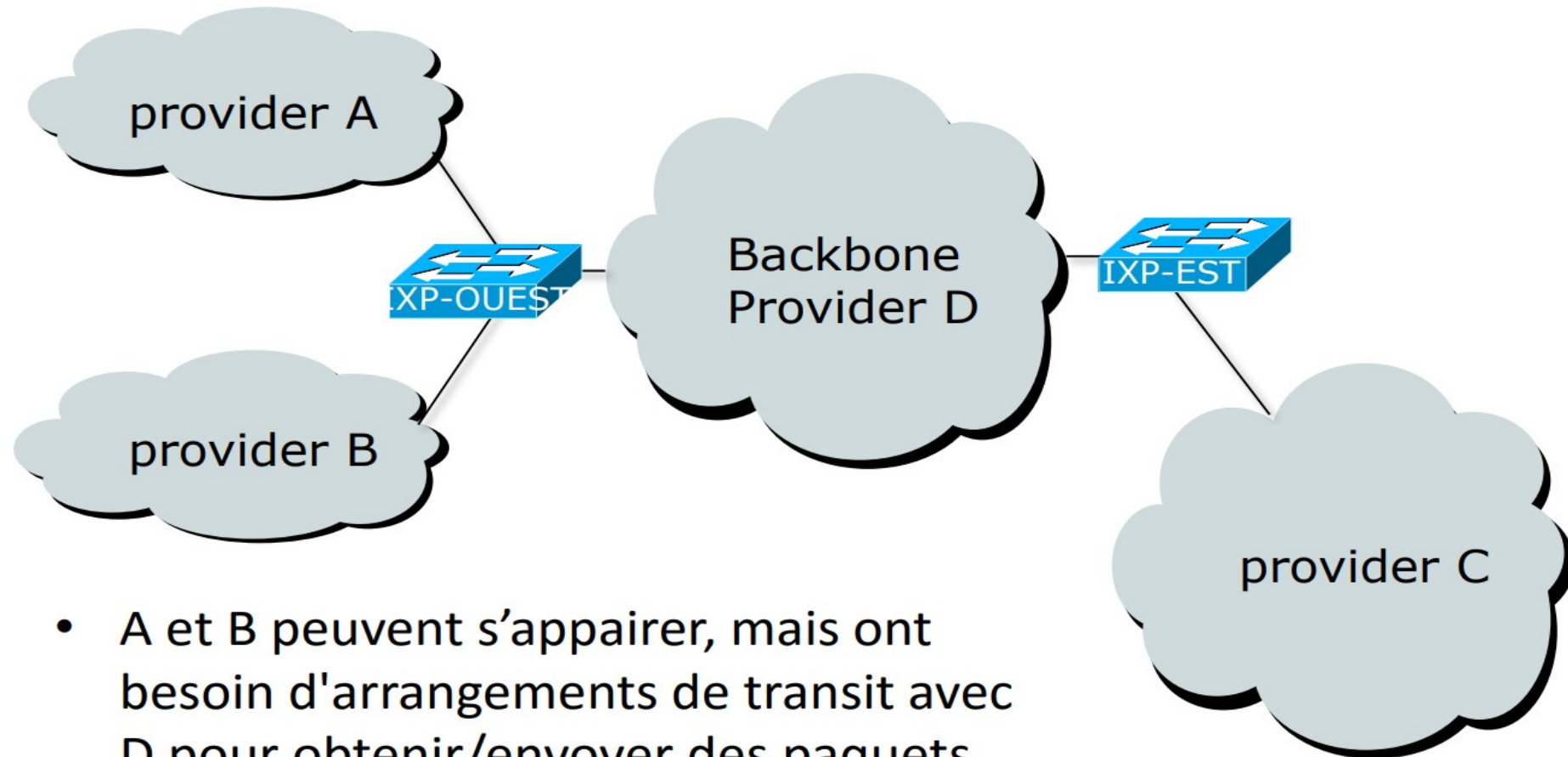
Default Free Zone

Default Free Zone

La Default Free Zone est constituée de routeurs Internet qui ont des informations de routage explicites sur le reste de l'Internet et n'ont donc pas besoin d'utiliser une route par défaut

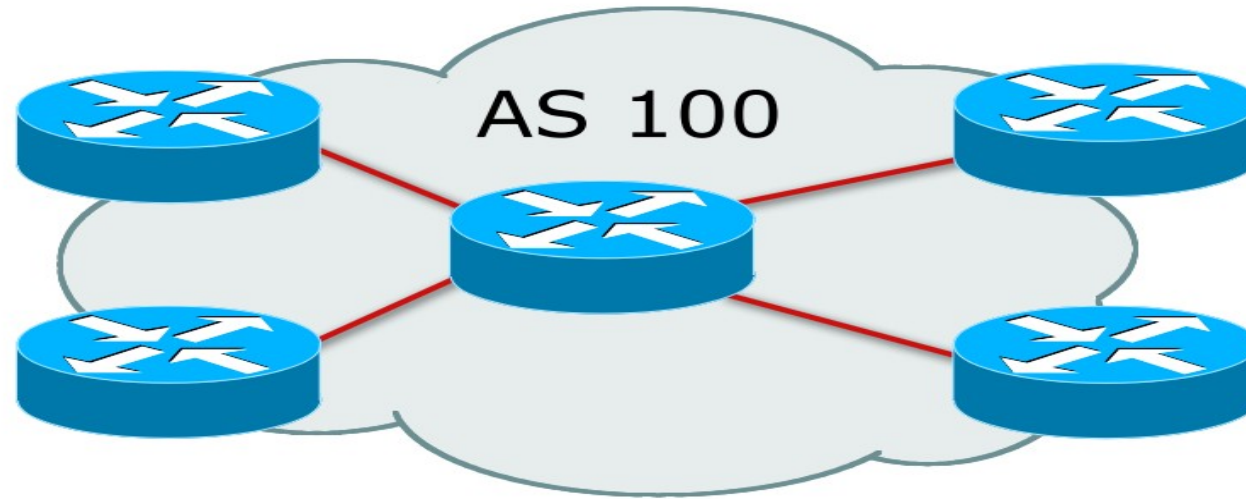
**NB: n'est pas liée à
l'emplacement d'un ISP dans
la hiérarchie**

Exemple de Peering et de Transit



- A et B peuvent s'appairer, mais ont besoin d'arrangements de transit avec D pour obtenir/envoyer des paquets de/vers C

Systeme autonome (AS)



- Regroupement de réseaux avec la même politique de routage
- Protocole de routage unique
- Habituellement en propriété, confiance et contrôle administratif unique
- Identifié par un entier 32 bits unique (ASN)

Numéro de système autonome (ASN)

- Deux plages
 - 0-65535 (original 16-bit range)
 - 65536-4294967295 (32-bit range – RFC4893)
- Usage:
 - 0 and 65535 (réservé)
 - 1-64495 (Internet public)
 - 64496-64511 (documentation – RFC5398)
 - 64512-65534 (usage privé uniquement)
 - 23456 (représente un AS32 bits si non géré)
 - 65536-65551 (documentation – RFC5398)
 - 65552-4294967295 (Internet public)
- Représentation de plage 32 bits spécifiée dans RFC5396
 - Définit “asplain” (format traditionnel) comme notation standard

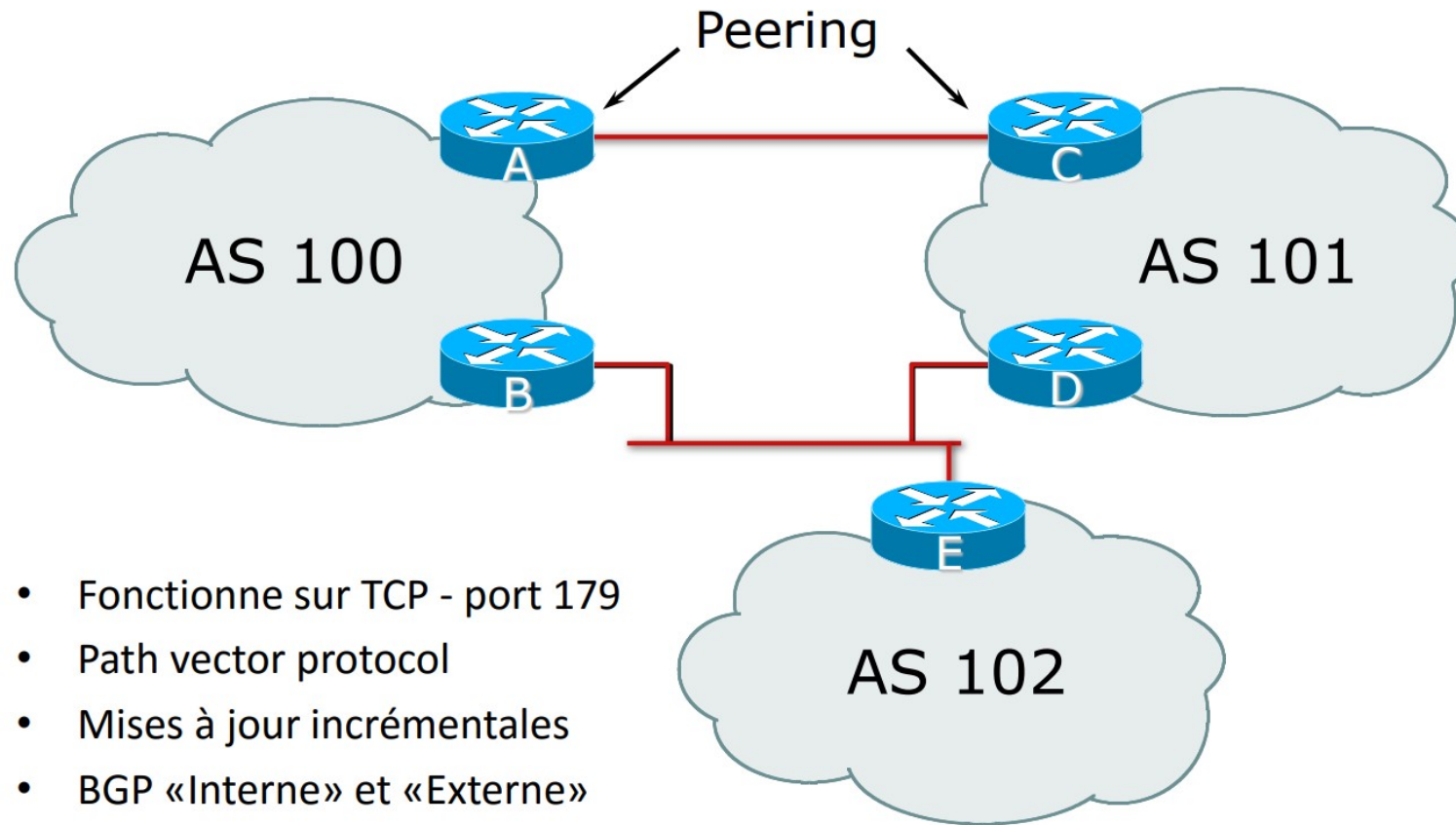
Numéro de système autonome (ASN): suite

- Les ASN sont distribués par les Registres Internet Régionaux
 - Ils sont également disponibles auprès des ISP en amont qui sont membres de l'un des RIR
- Les allocations ASN actuelles de 16-bit jusqu'à 61439 ont été faites pour les RIR
 - Environ 41200 sont visibles sur Internet
- Chaque RIR a également reçu un bloc d'ASN 32-bit
 - Sur les 2800 affectations, autour de 2400 sont visibles sur Internet
- Cf www.iana.org/assignments/as-numbers

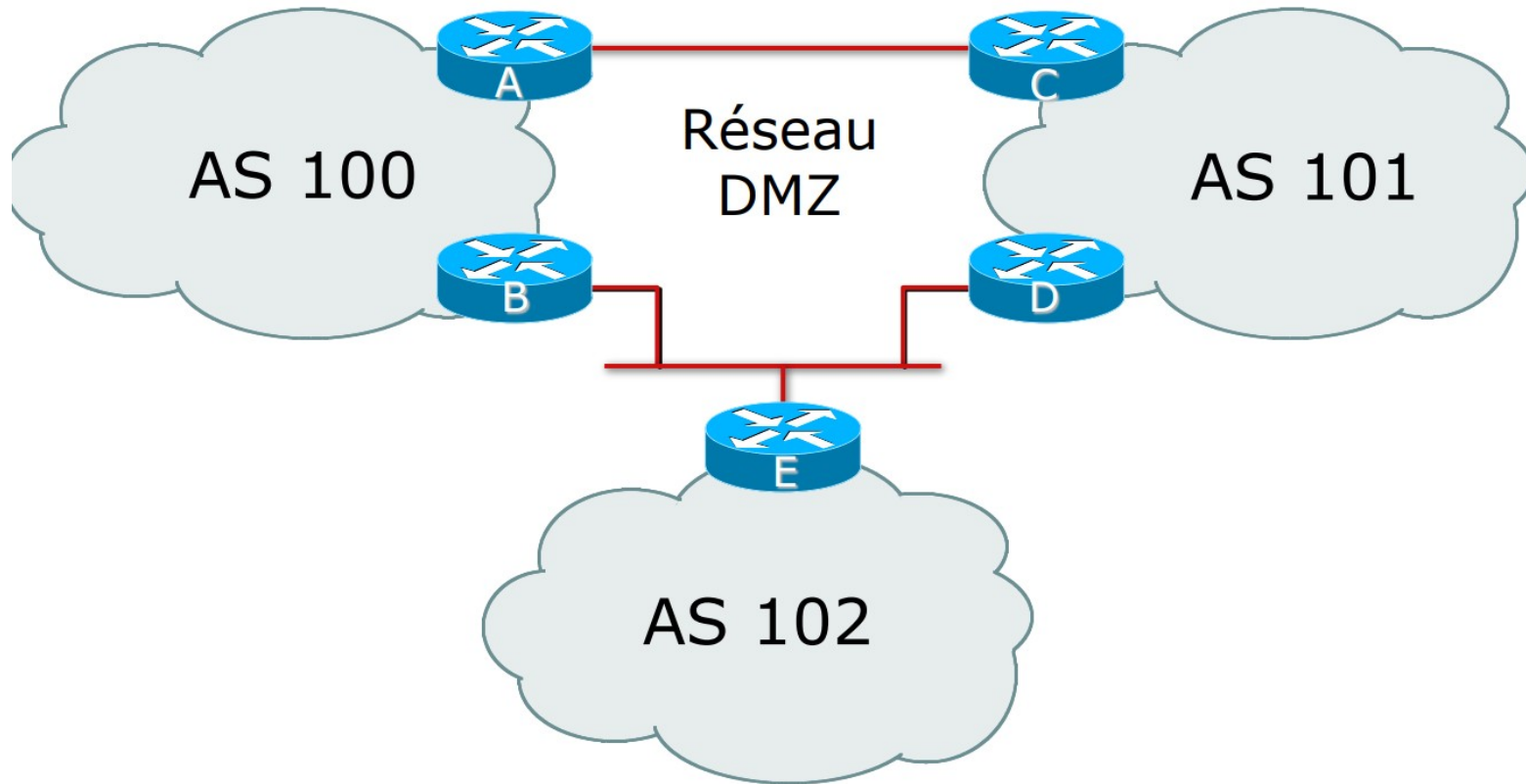
Configuration de base de BGP dans Cisco IOS

- Cette commande active BGP dans Cisco IOS:
`router bgp 100`
- Pour les ASN > 65535, le numéro d'AS peut être entré en notation "plain" ou "dot":
`router bgp 131076`
ou
`router bgp 2.4`
- IOS affiche les ASN en notation "plain" par défaut
 - La notation "Dot" est facultative:
`router bgp 2.4`
`bgp asnotation dot`

Notions de base BGP



Zone de démarcation (DMZ)



- DMZ est le lien ou réseau partagé entre les AS

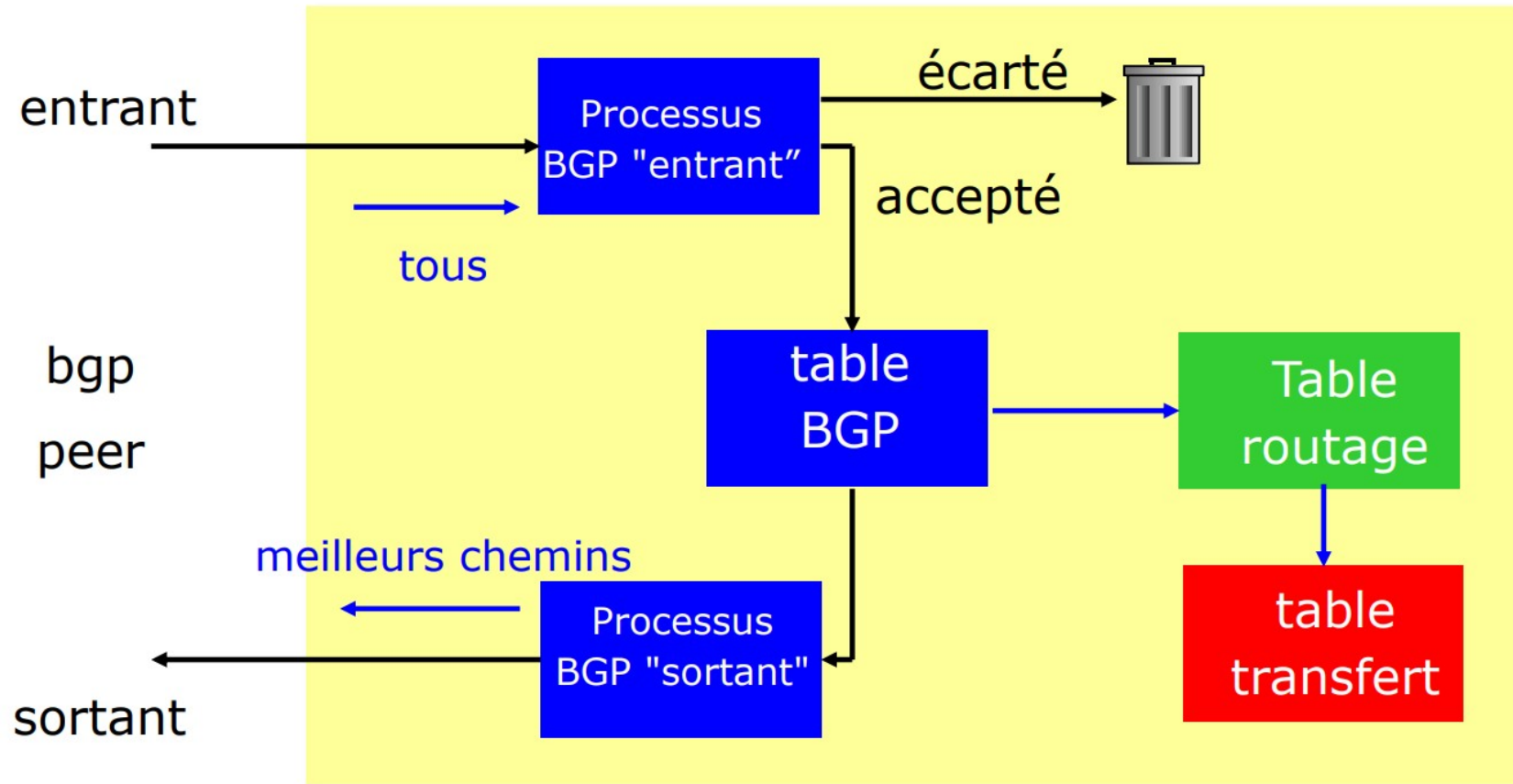
Opération générale BGP

- Apprend de multiples chemins par l'intermédiaire de speakers BGP internes et externes
- Choisit le meilleur chemin et l'installe dans la table de routage (RIB)
- Le meilleur chemin est envoyé aux voisins BGP externes
- Les stratégies sont appliquées en influant sur le choix du meilleur chemin

Construire la table de Transfert (Forwarding Table)

- Processus BGP “entrant”
 - reçoit les informations de chemin des pairs
 - les résultats de la sélection de chemin BGP placés dans la table BGP
 - “Meilleur chemin” marqué
- Processus BGP “sortant”
 - annonce le “meilleur chemin” aux pairs
- Meilleur chemin stocké dans la table de routage (RIB)
- Les meilleurs chemins dans le RIB sont ajoutés à la forwarding table (FIB) si:
 - préfixe et longueur de préfixe sont uniques
 - “distance administrative” la plus faible

Construire la table de Transfert (Forwarding Table)



Variante

Une façon de classer les protocoles de routage consiste à déterminer s'ils sont intérieurs ou extérieurs :

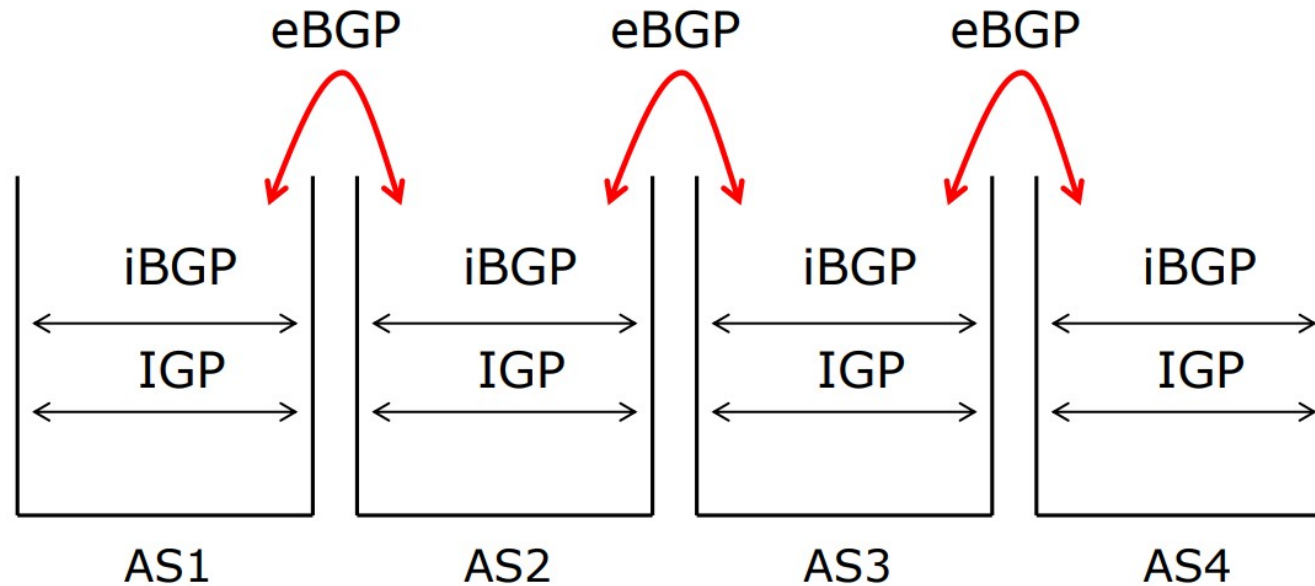
- Protocole de passerelle intérieure (**IGP**) : protocole de routage qui échange des informations de routage au sein d'un AS.
- Protocole de passerelle extérieure (**EGP**) : protocole de routage qui échange des informations de routage entre différents routeurs AS.

eBGP & iBGP

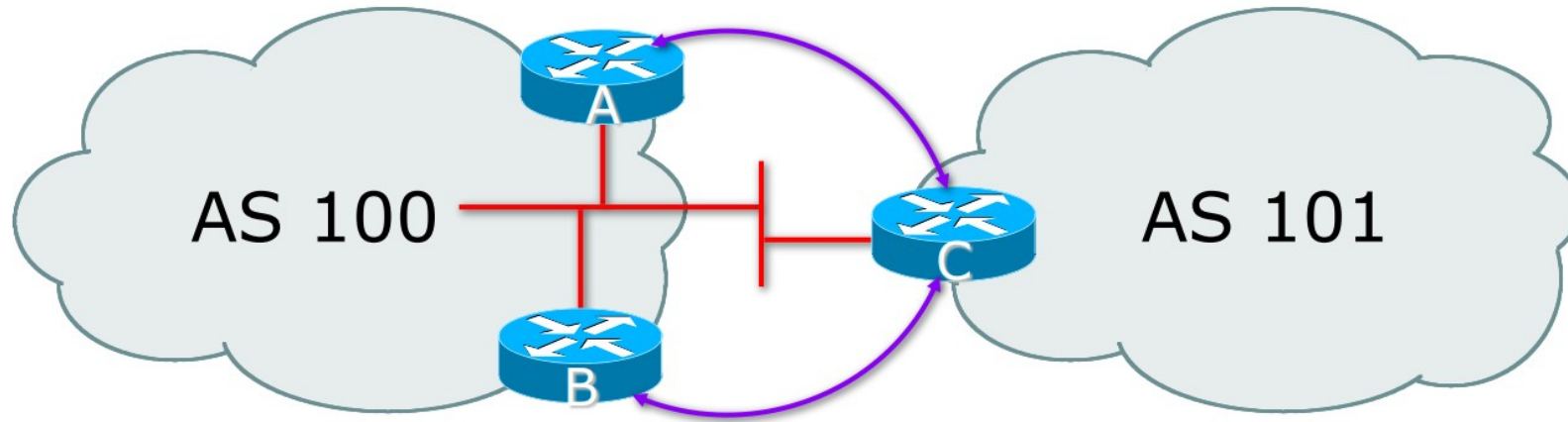
- BGP utilisé en interne (iBGP) et externe (eBGP)
- iBGP utilisé pour transporter
 - Certains / tous les préfixes Internet à travers le backbone ISP
 - les préfixes des clients ISP
- eBGP utilisé pour
 - Echanger des préfixes avec d'autres AS
 - Mettre en œuvre la politique de routage

Modèle BGP/IGP utilisé dans les réseaux ISP

- Représentation du modèle



External BGP Peering (eBGP)



- Entre speakers BGP de différents AS
- Doivent être directement connectés
- **Ne jamais** exécuter un IGP entre pairs eBGP

Configurer un BGP externe

Routeur A dans AS100

```
interface Ethernet 5/0
  ip address 102.102.10.2 255.255.255.240
!
router bgp 100
  network 100.100.8.0 mask 255.255.252.0
  neighbor 102.102.10.1 remote-as 101
  neighbor 102.102.10.1 prefix-list RouterC in
  neighbor 102.102.10.1 prefix-list RouterC out
!
```

adresse IP sur
interface Ethernet

ASN local

ASN distant

adresse ip de l'interface
Ethernet du routeur C

Filtres entrants
et sortants

Configurer un BGP externe

Routeur C dans AS101

```
interface ethernet 1/0/0
  ip address 102.102.10.1 255.255.255.240
!
router bgp 101
  network 100.100.64.0 mask 255.255.248.0
  neighbor 102.102.10.2 remote-as 100
  neighbor 102.102.10.2 prefix-list RouterA in
  neighbor 102.102.10.2 prefix-list RouterA out
!
```

adresse IP sur
interface Ethernet

ASN locale

ASN distant

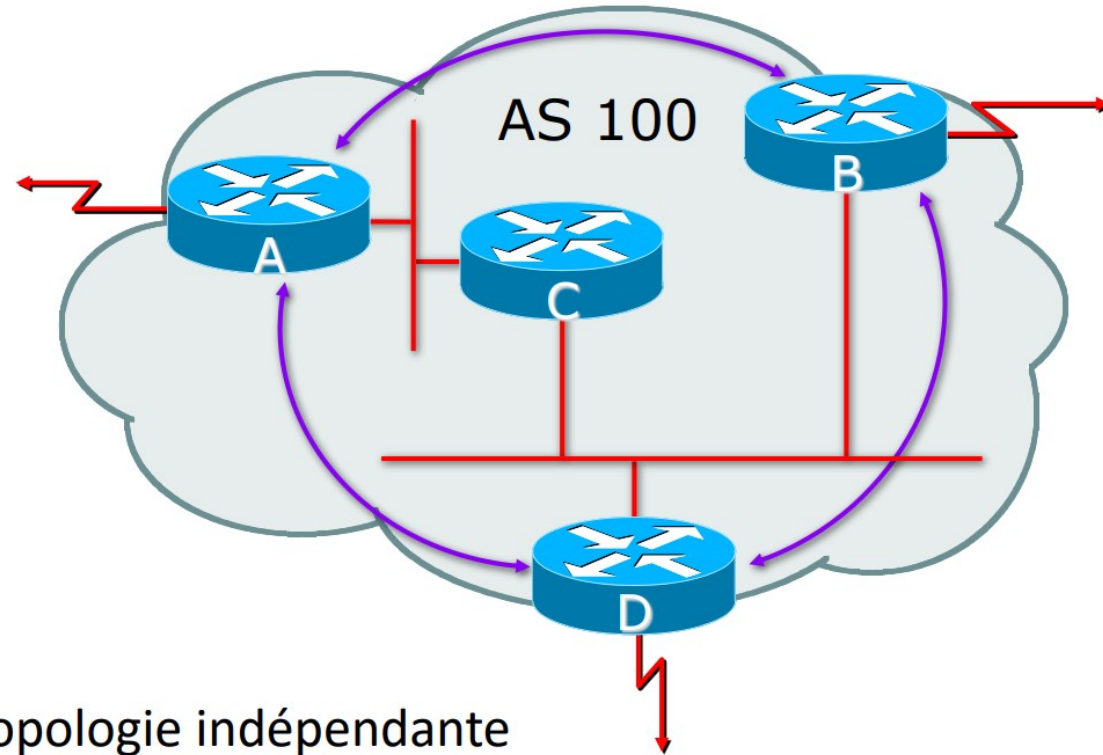
adresse ip de l'interface
Ethernet du routeur A

Filtres entrants
et sortants

BGP interne (iBGP)

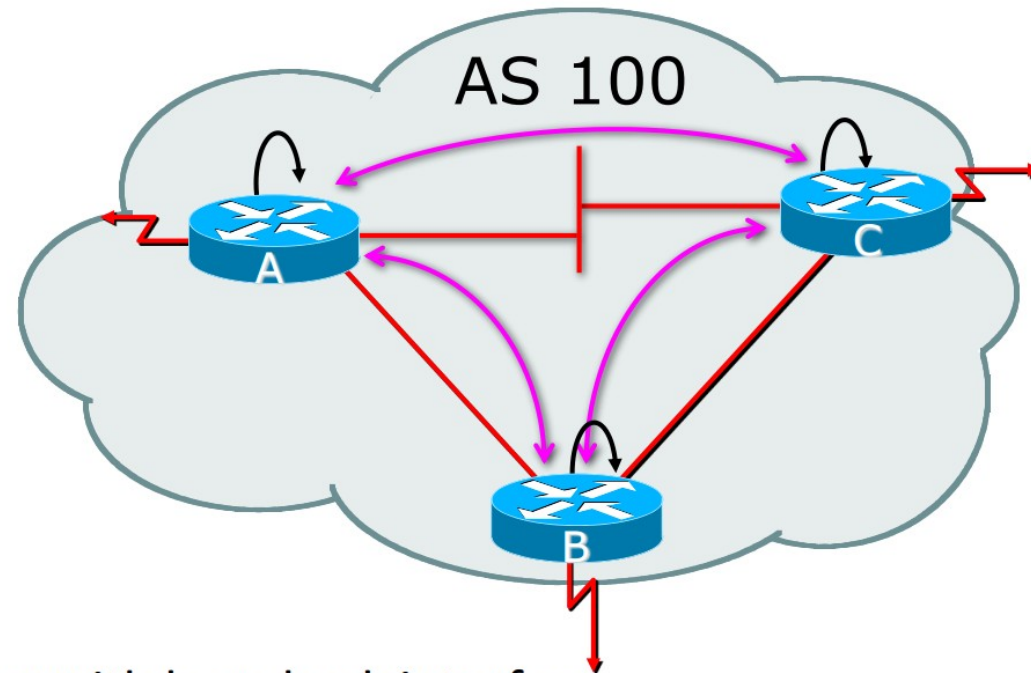
- BGP peer dans le même AS
- Pas requis d'être directement connecté
 - IGP prend soin de la connectivité des speakers inter-BGP
- Les speakers iBGP doivent être entièrement maillés:
 - Ils génèrent des réseaux connectés
 - Ils transmettent des préfixes appris de l'extérieur de l'ASN
 - Ils ne transmettent pas des préfixes appris d'autres voisins iBGP

Peering BGP interne (iBGP)



- Topologie indépendante
- Chaque speaker iBGP doit peering avec chacun des autres routeurs iBGP dans l'AS

Peering entre interfaces de Loopback



- Peer with loop-back interface
 - L'interface de loop-back ne tombe pas - jamais!
- Ne laissez pas une session iBGP dépendre de l'état d'une interface unique ou de la topologie physique

Configurer un BGP interne

Routeur A dans AS100

```
interface loopback 0
  ip address 105.3.7.1 255.255.255.255
!
router bgp 100
  network 100.100.1.0
  neighbor 105.3.7.2 remote-as 100
  neighbor 105.3.7.2 update-source loopback0
  neighbor 105.3.7.3 remote-as 100
  neighbor 105.3.7.3 update-source loopback0
!
```

adresse IP sur
interface loopback

ASN local

ASN local

adresse ip de l'interface
loopback du routeur B

Configurer un BGP interne

Routeur B dans AS100

```
interface loopback 0
 ip address 105.3.7.2 255.255.255.255
!
router bgp 100
 network 100.100.1.0
 neighbor 105.3.7.1 remote-as 100
 neighbor 105.3.7.1 update-source loopback0
 neighbor 105.3.7.3 remote-as 100
 neighbor 105.3.7.3 update-source loopback0
!
```

adresse IP sur
interface loopback

ASN local

ASN local

adresse ip de l'interface
loopback du routeur A

Insérer des préfixes dans BGP

Deux façons d'insérer des préfixes dans BGP

- Commande

redistribute static

- Commande

Network x.x.x.x

Insérer des préfixes dans BGP – redistribute static

- Exemple de configuration:

```
router bgp 100
  redistribute static
  ip route 102.10.32.0 255.255.254.0 serial0
```

- Une route statique doit exister avant que la commande « redistribute » fonctionne
- Force l'origine à être “incomplète”
- Soins requis!

Insérer des préfixes dans BGP – redistribute static

Soins requis avec redistribute!

- `redistribute<routing-protocol>` signifie que tout ce qui est dans le `<routing-protocol>` sera transféré dans BGP
- Ne s'ajustera pas s'il n'est pas contrôlé
- Préférable d'éviter autant que possible
- **redistribute** normalement utilisé avec des “route-maps” et sous contrôle administratif serré

Insérer des préfixes dans BGP – commande network

- Exemple de configuration

```
router bgp 100
  network 102.10.32.0 mask 255.255.254.0
  ip route 102.10.32.0 255.255.254.0 serial0
```

- Un itinéraire correspondant doit exister dans la table de routage avant que le réseau soit annoncé
- Force l'origine à être “IGP”

Configurer une agrégation

Trois façons de configurer l'agrégation

- `redistribute static`
- `adresse-agrégats`
- `network`

Configurer une agrégation

- Exemple de configuration:

```
router bgp 100
```

```
 redistribute static
```

```
 ip route 102.10.0.0 255.255.0.0 null0 250
```

- une route statique vers “null0” est appelée une route « pull up »
 - Paquets envoyés ici uniquement s'il n'y a plus de correspondance spécifique dans la table de routage
 - distance de 250 assure que ceci est le dernier recours statique
 - soins requis - voir plus haut!

Configurer une agrégation – Commande network

- Exemple de configuration

```
router bgp 100
  network 102.10.0.0 mask 255.255.0.0
  ip route 102.10.0.0 255.255.0.0 null0 250
```

- Un itinéraire correspondant doit exister dans la table de routage avant que le réseau soit annoncé
- Moyen le plus simple et le meilleur de générer un agrégat

Configurer une agrégation – commande `aggregate-address`

- Exemple de configuration:

```
routeur bgp 100
```

```
network 102.10.32.0 mask 255.255.252.0
```

```
aggregate-address 102.10.0.0 255.255.0.0 [summary-only]
```

- Nécessite un préfixe plus spécifique dans le tableau BGP avant que l'agrégat soit annoncé
- Mot-clé uniquement de résumé (summary-only keyword)
 - Mot-clé facultatif qui garantit que seul le résumé est annoncé si un préfixe plus spécifique existe dans la table de routage

Résumé BGP neighbour status

```
Router6>sh ip bgp sum
BGP router identifier 10.0.15.246, local AS number 10
BGP table version is 16, main routing table version 16
7 network entries using 819 bytes of memory
14 path entries using 728 bytes of memory
2/1 BGP path/bestpath attribute entries using 248 bytes of memory
0 BGP route-map cache entries using 0 bytes of memory
0 BGP filter-list cache entries using 0 bytes of memory
BGP using 1795 total bytes of memory
BGP activity 7/0 prefixes, 14/0 paths, scan interval 60 secs
```

Neighbor	V	AS	MsgRcvd	MsgSent	TblVer	InQ	OutQ	Up/Down	State/PfxRcd
10.0.15.241	4	10	9	8	16	0	0	00:04:47	2
10.0.15.242	4	10	6	5	16	0	0	00:01:43	2
10.0.15.243	4	10	9	8	16	0	0	00:04:49	2
...									

Version BGP

Mises à jour
envoyées et
reçues

Mises à jour en attente

Résumé Table BGP

```
Router6>sh ip bgp
```

```
BGP table version is 30, local router ID is 10.0.15.246
```

```
Status codes: s suppressed, d damped, h history, * valid, > best, i -  
internal,
```

```
                r RIB-failure, S Stale
```

```
Origin codes: i - IGP, e - EGP, ? - incomplete
```

Network	Next Hop	Metric	LocPrf	Weight	Path
*>i10.0.0.0/26	10.0.15.241	0	100	0	i
*>i10.0.0.64/26	10.0.15.242	0	100	0	i
*>i10.0.0.128/26	10.0.15.243	0	100	0	i
*>i10.0.0.192/26	10.0.15.244	0	100	0	i
*>i10.0.1.0/26	10.0.15.245	0	100	0	i
*> 10.0.1.64/26	0.0.0.0	0		32768	i
*>i10.0.1.128/26	10.0.15.247	0	100	0	i
*>i10.0.1.192/26	10.0.15.248	0	100	0	i
...					

Le protocole VPN-MPLS et TE

Introduction

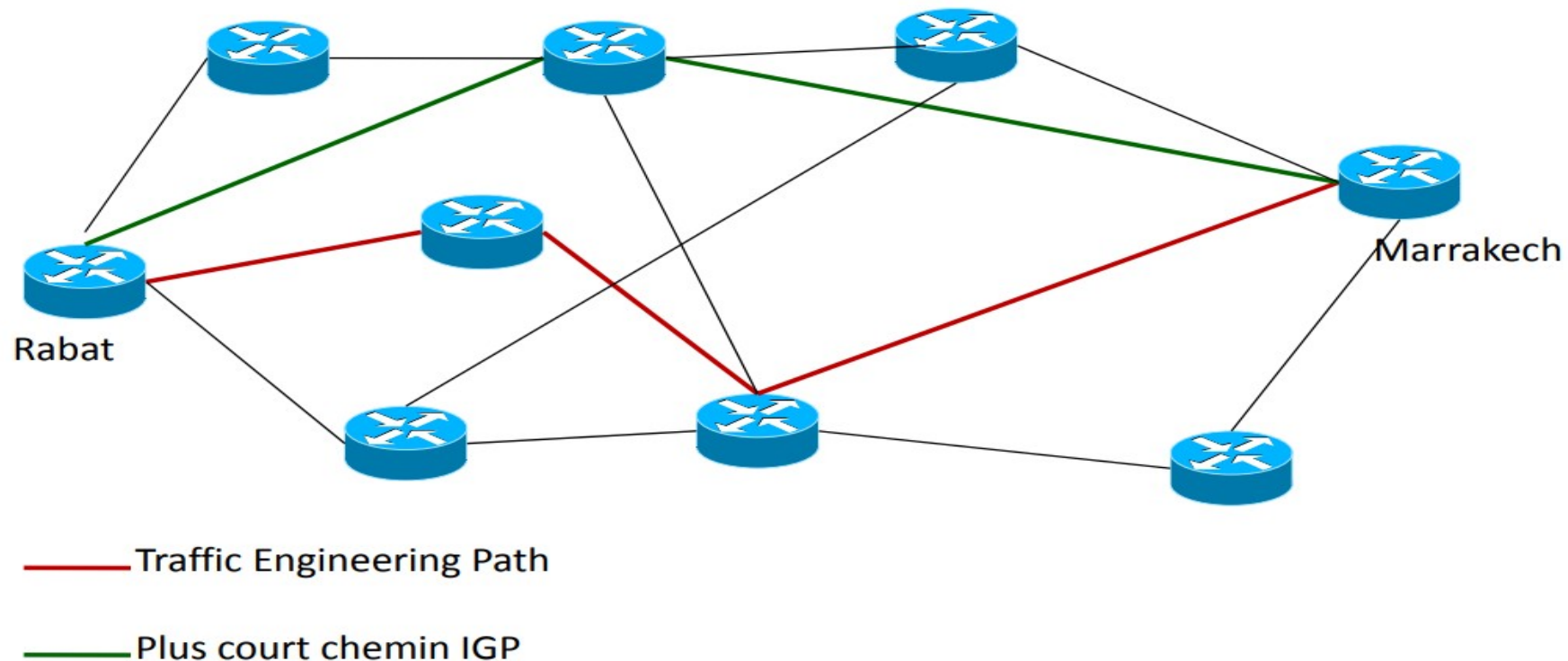
- Les VPN/MPLS sont essentiellement implémentés chez les opérateurs afin de fournir des services à leurs clients. Les opérateurs utilisent leur backbone sur MPLS pour créer des VPN, par conséquent le réseau MPLS des opérateurs se trouve partagé ou mutualisé avec d'autre client. Du point de vue du client, il a l'impression de bénéficier d'un réseau qui lui est entièrement dédié. C'est-à-dire qu'il a l'impression d'être le seul à utiliser les ressources que l'opérateur lui met à disposition. Ceci est dû à l'étanchéité des VPN/MPLS qui distingue bien les VPN de chaque client et tous ces mécanismes demeurent transparents pour les clients. Finalement, les deux parties sont gagnantes car les clients ont un véritable service IP qui leur offre des VPN fiables à des prix plus intéressants que s'ils devaient créer eux-mêmes leur VPN de couche 2. Les opérateurs eux aussi réduisent leurs coûts du fait de la mutualisation de leurs équipements. Voici une représentation des VPN/MPLS

Les applications MPLS

- ❑ **L'ingénierie de trafic** : Les administrateurs de réseau peuvent mettre en oeuvre des politiques visant à assurer une distribution optimale du trafic et à améliorer l'utilisation globale du réseau.

Les applications MPLS: MPLS-TE

- Router les données à travers le réseau suivant une vision de management de la disponibilité des ressources, et la qualité de service, mais toujours à base de plus court chemin comme le cas du routage IP.



Les applications MPLS:

Objectifs de l'ingénierie de trafic: RFC 270

❑ Vise à optimiser les réseaux opérationnels et d'augmenter leurs performances en terme de ressource et de trafic.

❑ Ces objectif majeurs (RFC:2702):

- Objectif orienté trafic:

1. Minimiser les pertes.
2. Minimiser le délai.
3. Maximiser les transferts « throughput ».
4. Renforcer SLA.

-Objectif orienté ressource:

1. Optimiser l'utilisation du réseau.

Les applications MPLS: La bande passante garantie

- ❑ permet aux fournisseurs de services d'allouer des largeurs de bande passante et des canaux garantis.
- ❑ Ce qui permet également la comptabilité des ressources QoS (qualité de service) de manière à organiser le trafic 'prioritaire' et 'au mieux', tels que la voix et les données.

Les applications MPLS: Le Re-routage rapide

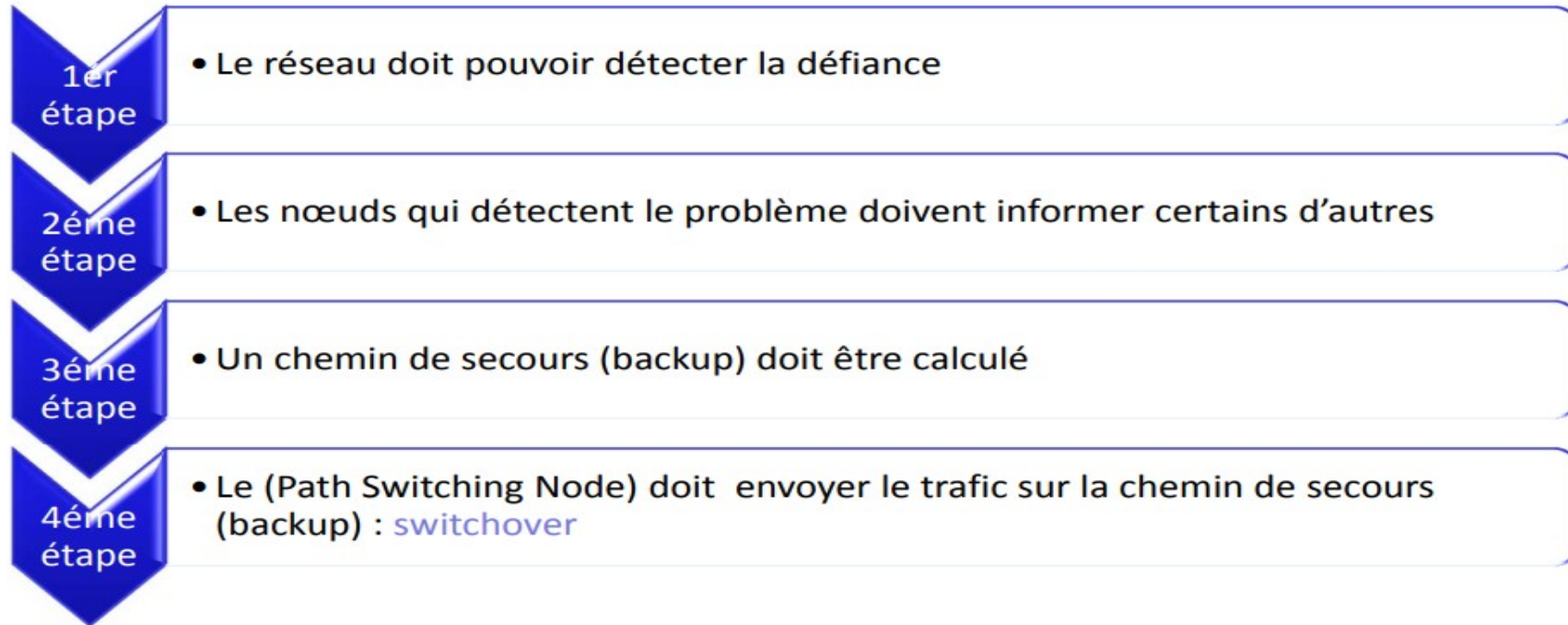
- ❑ **permet** une reprise très rapide après la défaillance d'une liaison ou d'un nœud.
- ❑ Une telle rapidité de reprise empêche l'interruption des applications utilisateur ainsi que toute perte de données.

Problématique :

Les pannes des éléments du réseau peuvent être d'origines diverses.

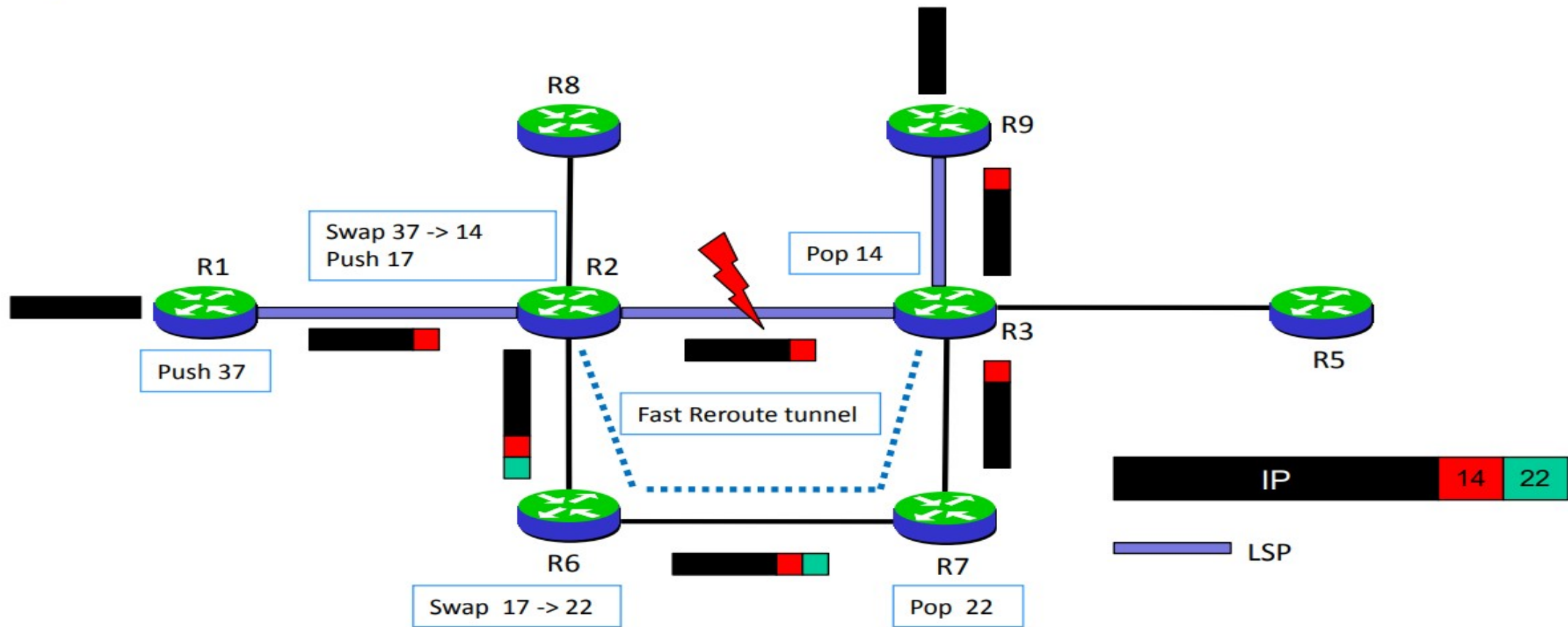
Type de la panne	Distribution
Software upgrade/configuration	22%
Hardware upgrade/configuration	9%
Software Failure	25%
Hardware Failure	14%
Link Failure	20%
Power Outage	1%
Other	9%

Problématique :



Temps de rupture = Duré de détection d'erreur
+ Durée de notification d'erreur
+ Durée nécessaire pour calculer le chemin de backup
+ Durée de switchover.

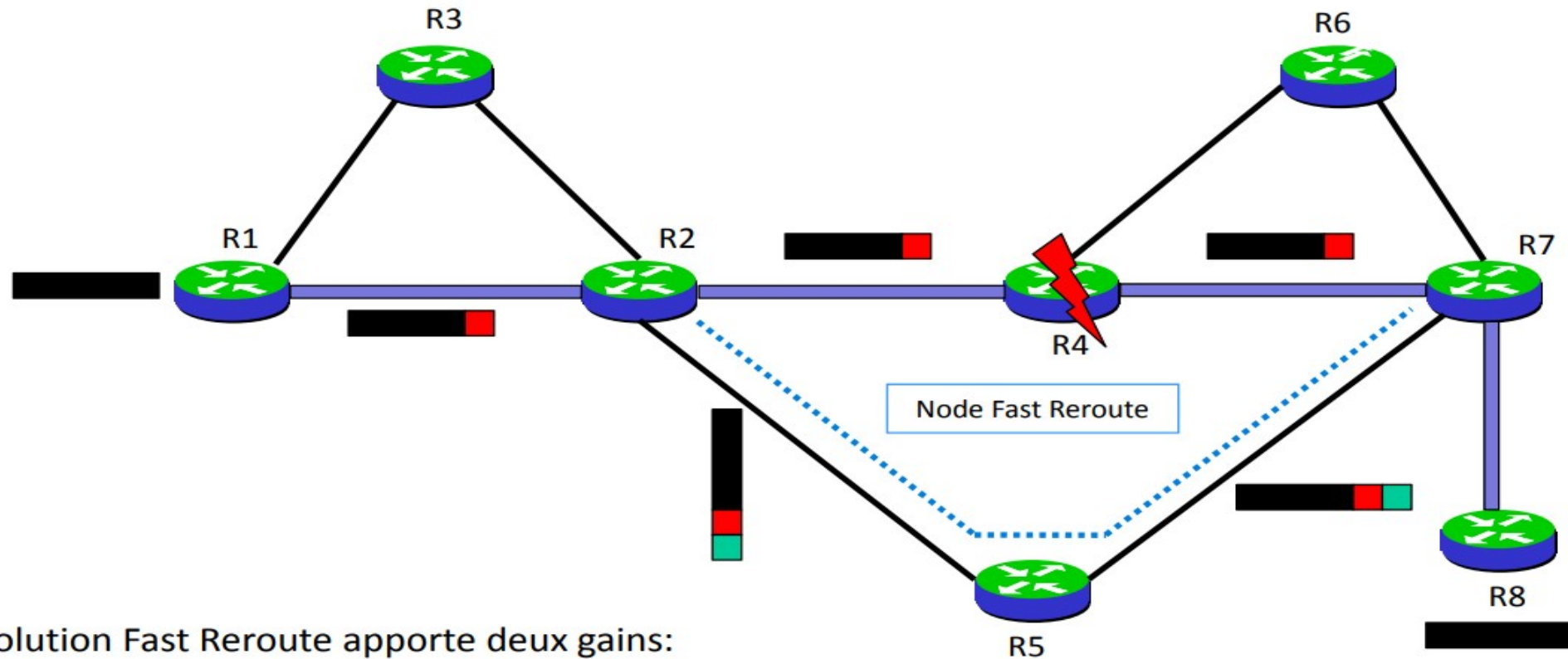
Etapes d'un Fast Reroute de lien



On crée un tunnel entre R2 et R3 pour rerouter le trafic.

Le temps de recouvrement est rapide puisque le reroutage se fait par R2 et pas par R1

Etapes d'un Fast Reroute de nœud



La solution Fast Reroute apporte deux gains:

- Une fiabilité accrue pour les services IP: on est protégé contre les ruptures des services.
- Une Scalabilité importante inhérente au design du réseau, puisque le chemin de backup est en fonction du lien et pas en fonction des LSPs qui traversent ce lien

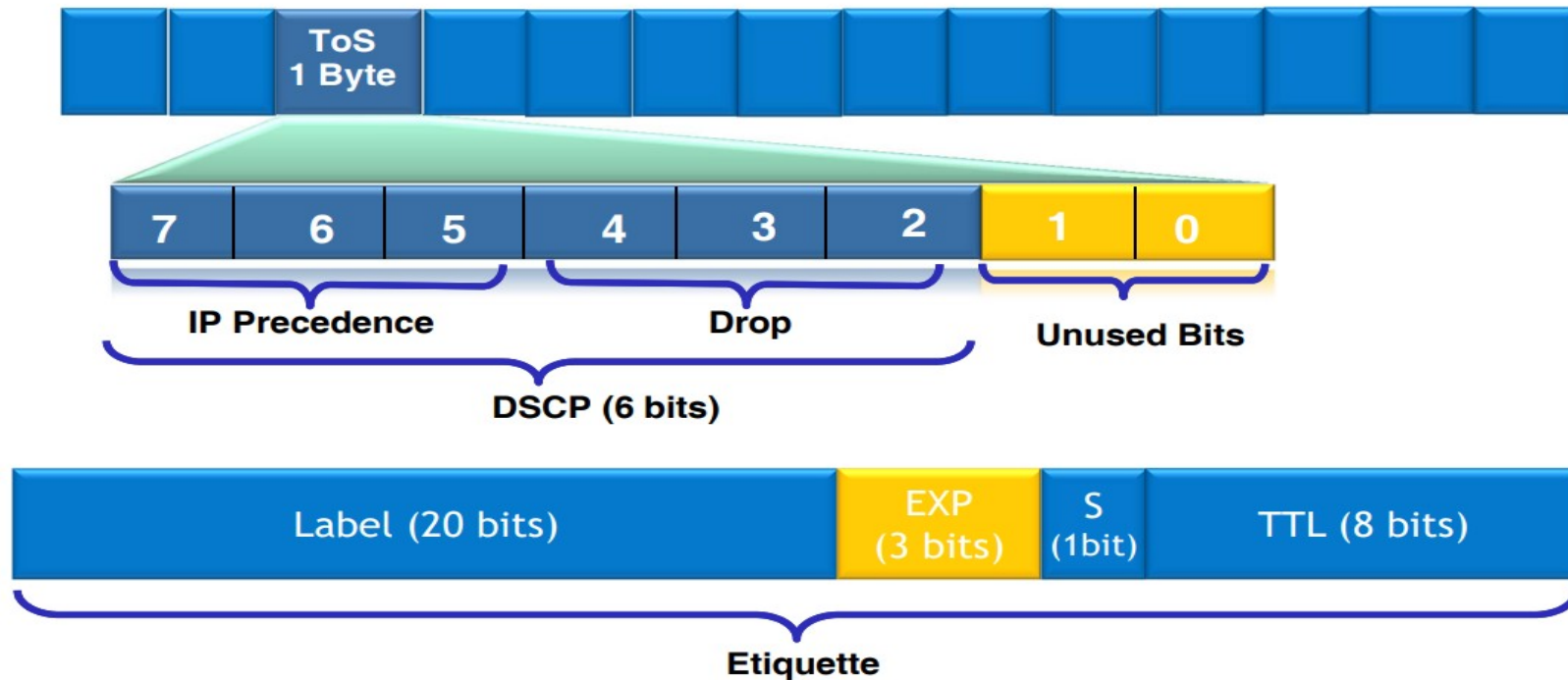
Les applications MPLS: La fonction Classe de service (CoS)

- ❑ La fonction Classe de service (CoS) MPLS assure que le trafic important est traité avec la priorité adéquate sur le réseau et que les exigences de latence sont respectées.
- ❑ Les mécanismes QoS IP peuvent être mis en oeuvre de façon transparente dans un environnement MPLS.

MPLS-DiffServ

- ❑ Le champ DSCP a 6 bits tandis que le champ EXP (experimental field) des étiquettes a seulement 3 bits.

Couche 3 (IPV4)



MPLS-DiffServ:

- ❑ Deux solutions ont été proposées pour résoudre ces deux problèmes.
 - ❑ **E-LSP (EXP-Inferred-PSC LSPs) .**
 - ❑ **L-LSP (Label-Only-Inferred-PSC LSPs) .**

MPLS-DiffServ : (E-LSP, L-LSP)

Un LSP peut appartenir à plusieurs classes de service. Mais l'activation de ces PHB dans chaque LSR est faite manuellement : ni LDP, ni RSVP-TE n'acheminent actuellement ces éléments.

Dans un L-LSP, le label sert à identifier une FEC et ses priorités de traitement. Dans un L-LSP, le champ EXP est utilisé uniquement pour le niveau de perte.

E-LSP



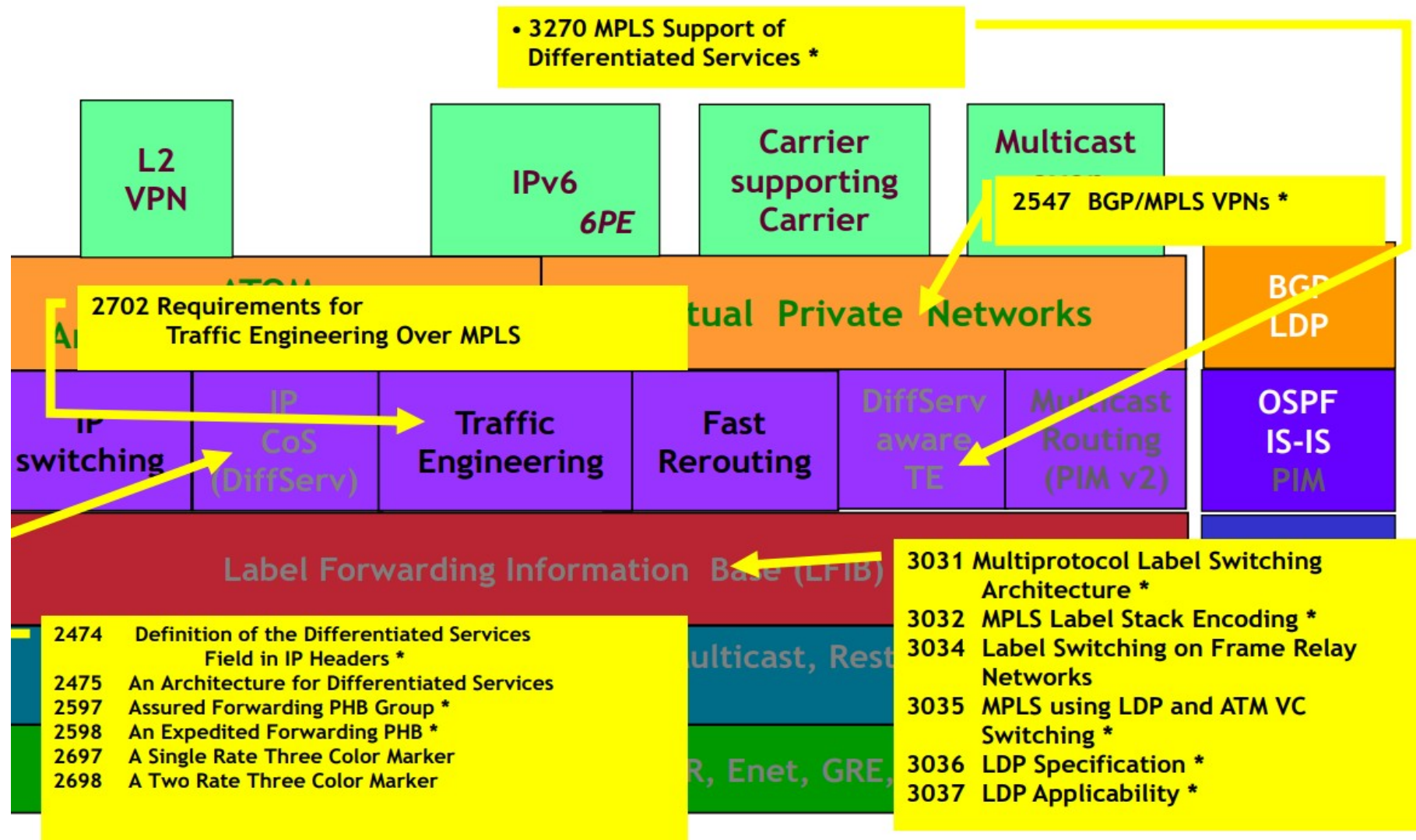
L-LSP



Les services MPLS

- ❑ MPLS trouve de nouveaux domaines d'applications pour s'approcher de plus en plus des couches applicatives.
- ❑ MPLS permet aux opérateurs de mettre à la disposition des clients professionnels un certain nombre de services à valeur ajoutée pouvant se greffer au dessus de MPLS.

MPLS et Standards



Les applications MPLS: Les réseaux privés virtuels MPLS

- ❑ Les réseaux privés virtuels MPLS simplifient considérablement le déploiement des services par rapport aux VPN IP traditionnels.

Le pourquoi du VPN

❑ Problème posé

- ❑ Comment permettre à plusieurs réseaux indépendants de cohabiter sur une même infrastructure (FAI) ?
 - ❑ *Ces réseaux doivent pouvoir utiliser le même adressage privé*
 - ❑ *Les trafics issus des différents réseaux doivent être isolés*

❑ Solutions proposées

- ❑ Solutions classiques
 - ❑ *Utilisation de lignes spécialisées*
 - ❑ *Utilisation de tunnels IP*
 - ❑ *Utilisation de réseaux Frame Relay ou ATM*
- ❑ Solutions alternatives
 - ❑ *Utilisation simple de MPLS*
 - ❑ *Utilisation améliorée de MPLS*

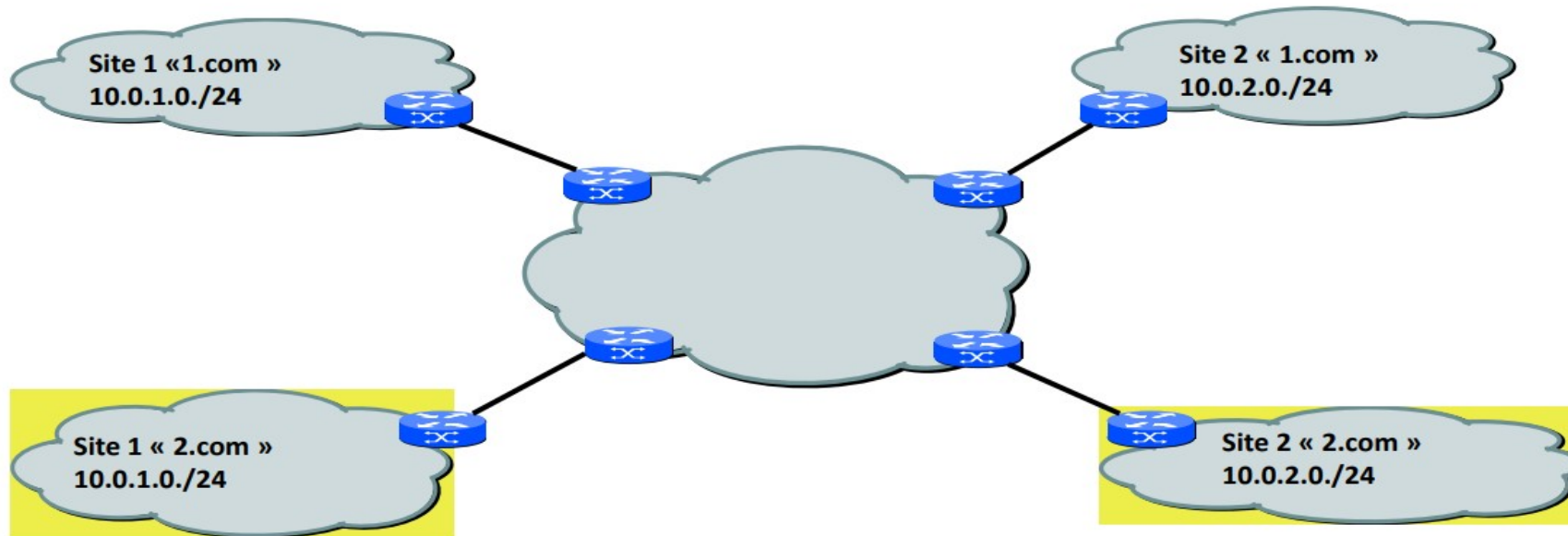
Le pourquoi du VPN

❑ Généralités

- ❑ Les réseaux publics (i.e. Internet) en tant qu'alternative aux liaisons spécialisées
 - ❑ Solutions meilleure marché
 - ❑ Solutions non-sécurisées par nature
- ❑ Usage de « tunnels »
 - ❑ Solutions permettant d'étendre le réseau interne
 - ❑ Solutions permettant d'encapsuler les données à transporter
 - ❑ Solutions généralement associées à une sécurisation des échanges (authentification, chiffrement)

Exemple d'architecture VPN

- ❑ **Comment interconnecter dans deux VPN distincts**
 - ❑ Les réseaux des sites « 1.com » et « 2.com »
 - ❑ En supposant que les deux réseaux utilisent le même adressage privé « 10.0.X.0/24 » pour le site X



Solutions classiques (1)

❑ Utilisation de Lignes Spécialisées

❑ Avantages

- ❑ *Qualité de service a priori de très bonne qualité*

❑ Inconvénients

- ❑ *Coût pouvant être très élevé*

- ❑ Même en cas de trafic inter-sites faible

- ❑ *Nombre de lignes à louer potentiellement important*

- ❑ $N.(N-1)$ pour une architecture à maillage complet

- ❑ *Solution peu évolutive*

- ❑ Pour chaque nouveau site, au moins une nouvelle LS doit être mise en place

Solutions classiques (1)

❑ Utilisation de Tunnels IP

❑ Avantages

1. *Coût moins important*

- ❑ Partage plus important des ressources du backbone

❑ Inconvénients

- ❑ *Qualité de Service non nécessairement garantie*

- ❑ *Nombre de tunnels à configurer potentiellement important*

- ❑ $N.(N-1)$ pour une architecture à maillage complet

- ❑ Solution évolutive

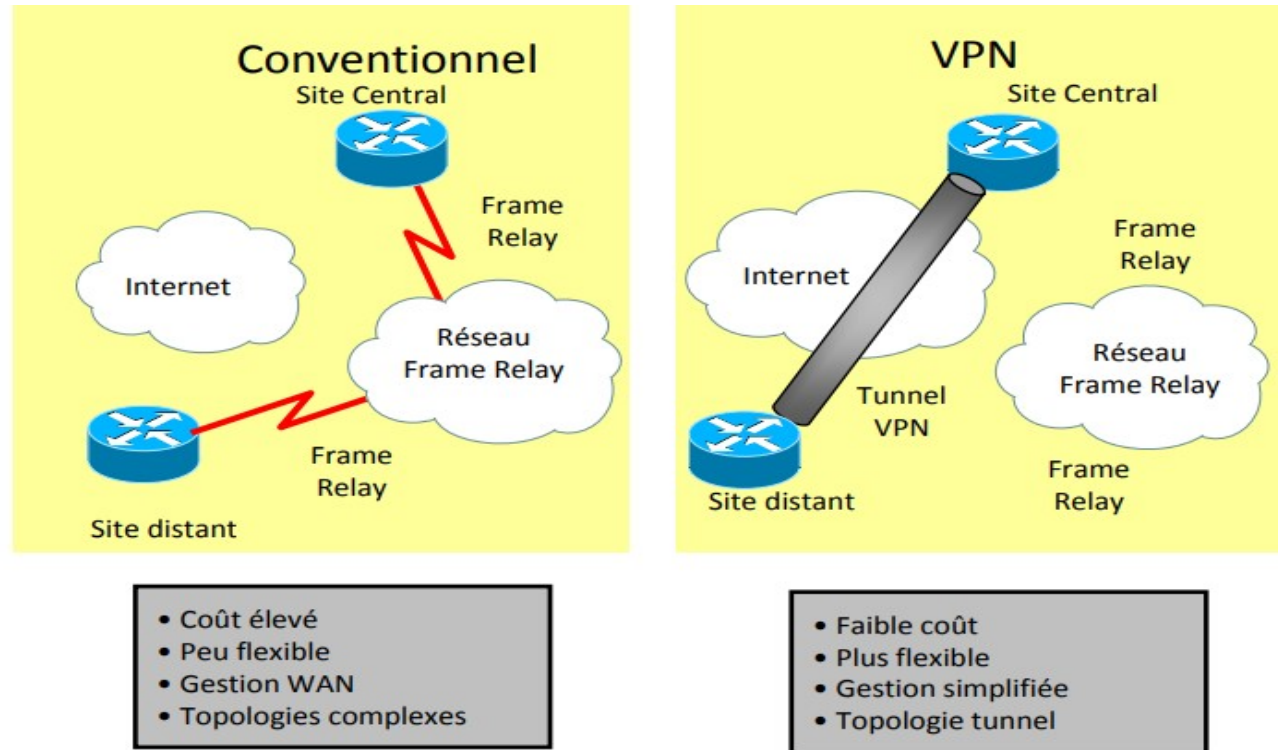
- ❑ *Sécurité dépendante du mécanisme de « tunneling » utilisé*

- ❑ GRE (Generic Routing Encapsulation), Ipsec, TLS,....

- ❑ *Configuration lourde*

- ❑ Principalement à la charge des clients

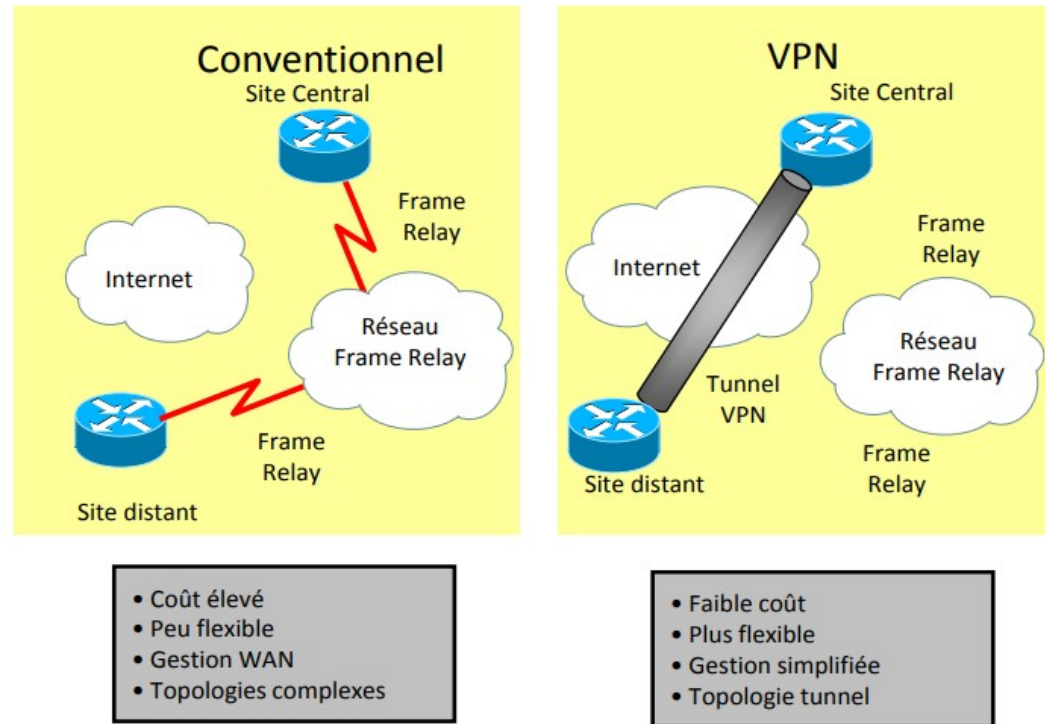
Présentation des VPN



- Les principaux avantages sont:

- Les VPNs amènent des coûts plus faibles que les réseaux privés.
- Les coûts de la connectivité LAN-LAN sont réduits de 20 à 40 pourcent par rapport à une ligne louée. .
- Les VPNs offrent plus de flexibilité et d'évolutivité que des architectures WAN classiques

Présentation des VPN



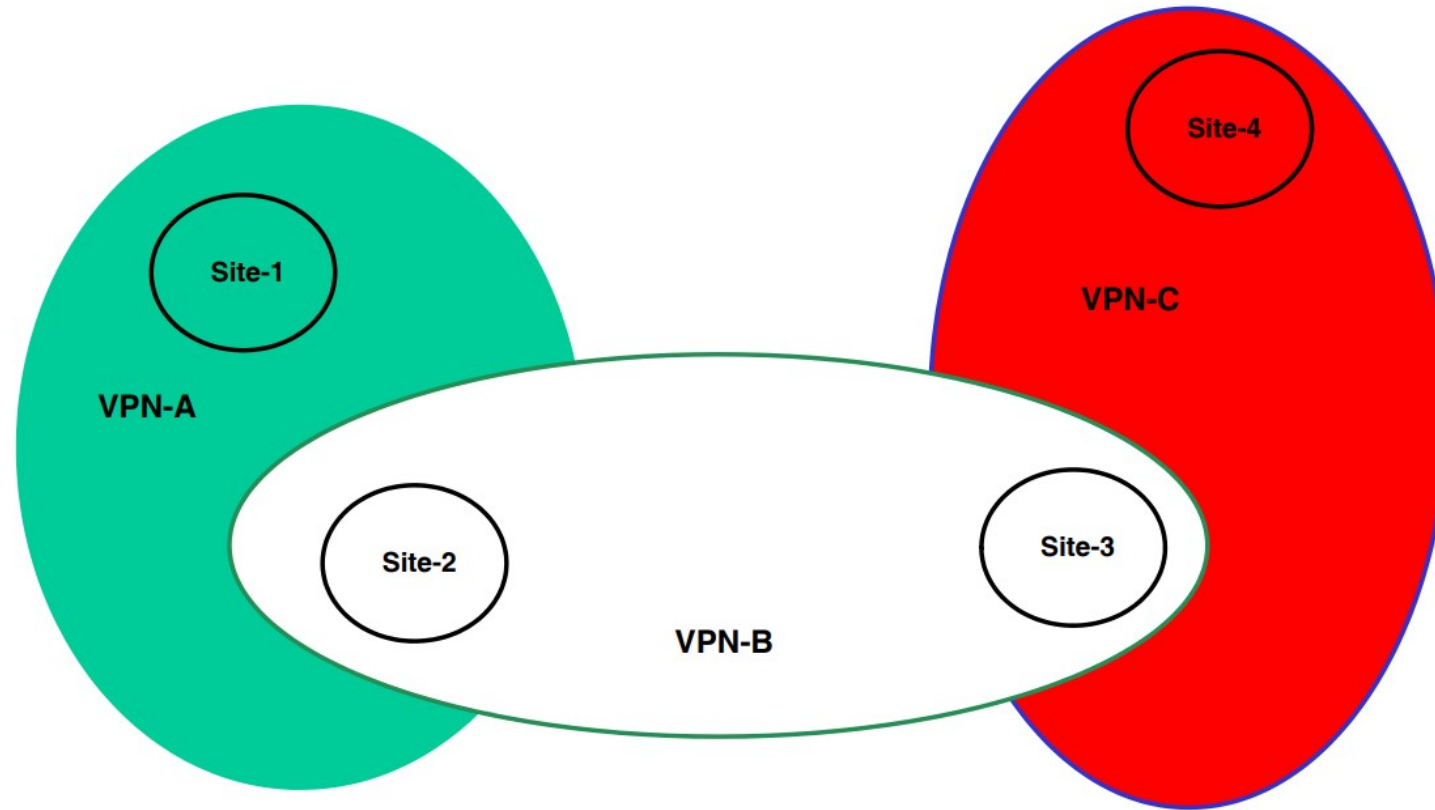
- Les principaux avantages sont:

- Les VPNs simplifient les tâches de gestion comparé à la l'exploitation de sa propre infrastructure de réseau.
- Les VPNs fournissent des topologies de réseaux avec tunnels qui réduisent les taches de gestion.
- Un backbone IP n'utilise pas les circuits virtuels permanents (PVCs) avec des protocoles orientés connexion tels ATM et FR

Le service VPN-IP

- ❑ Réseau Privé bâti sur une infrastructure mutualisée;
- ❑ Absence de lien physique de bout en bout;
- ❑ Confidentialité est assurée par la préservation des plans de routage et d'adressage;
- ❑ Le concept de réseau privé virtuel ou VPN n'est pas nouveau;
 - ❑ Utilisation des circuits virtuels;
 - ❑ Emulation des liens point à point pour le client;
 - ❑ Partage statistique de l'infrastructure du SP

Topologie VPN

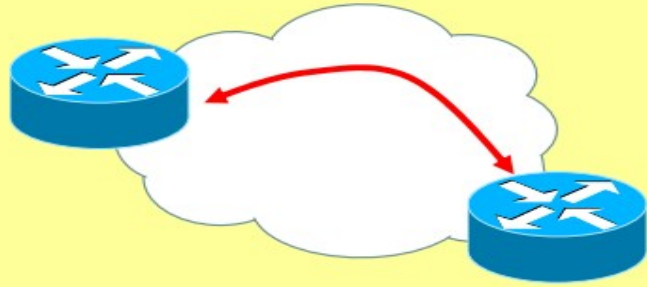


Les postes de base

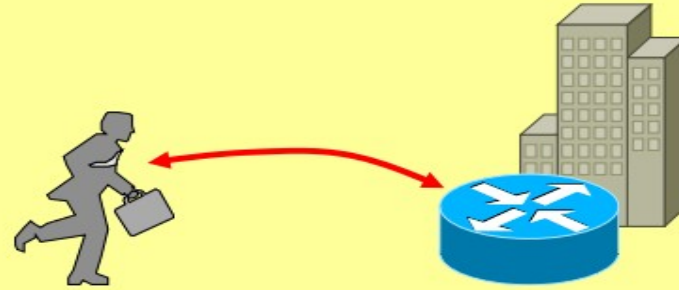
- ❑ Trois composantes fondamentales :
 - ❑ CE (Customer Edge);
 - ❑ PE (Provider Edge);
 - ❑ P (Provider Router);

Scénarios VPN

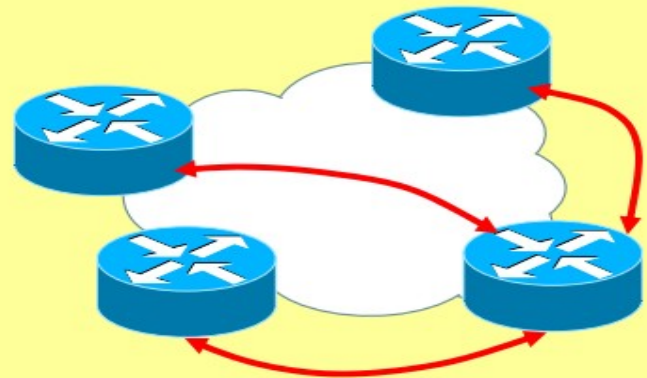
Routeur à Routeur



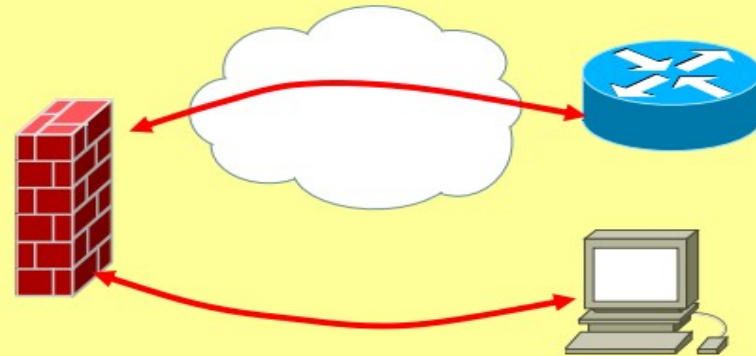
PC à Routeur/Concentrateur



Routeur à plusieurs routeurs



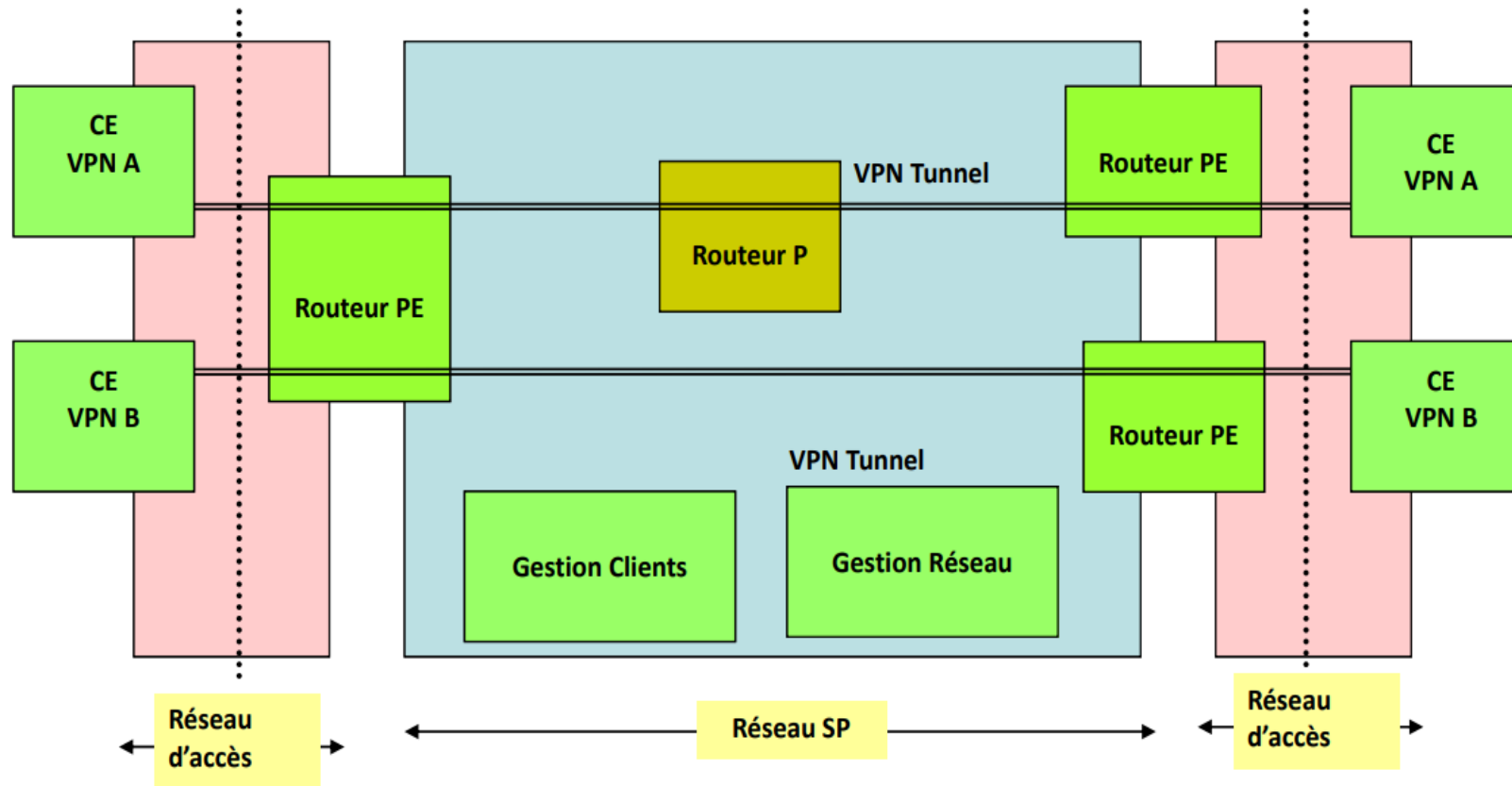
PC à Pare-Feu



Les modèles de référence

- ❑ Trois principaux modèles VPN :
 - ❑ CPE-Based VPN;
 - ❑ Network-based layer 3 IPVPN ;
 - ❑ Network-based layer 2 IPVPN ;
- ❑ La connexion CE-PE consiste en un circuit physique dédié :
 - ❑ un circuit logique (FR ou ATM)
 - ❑ un tunnel IP (utilisant par exemple IPsec, L2TP).
- ❑ Dans le backbone SP, les tunnels VPN sont normalement utilisés pour interconnecter les équipements PEs.

Le modèle CPE-Based VPN

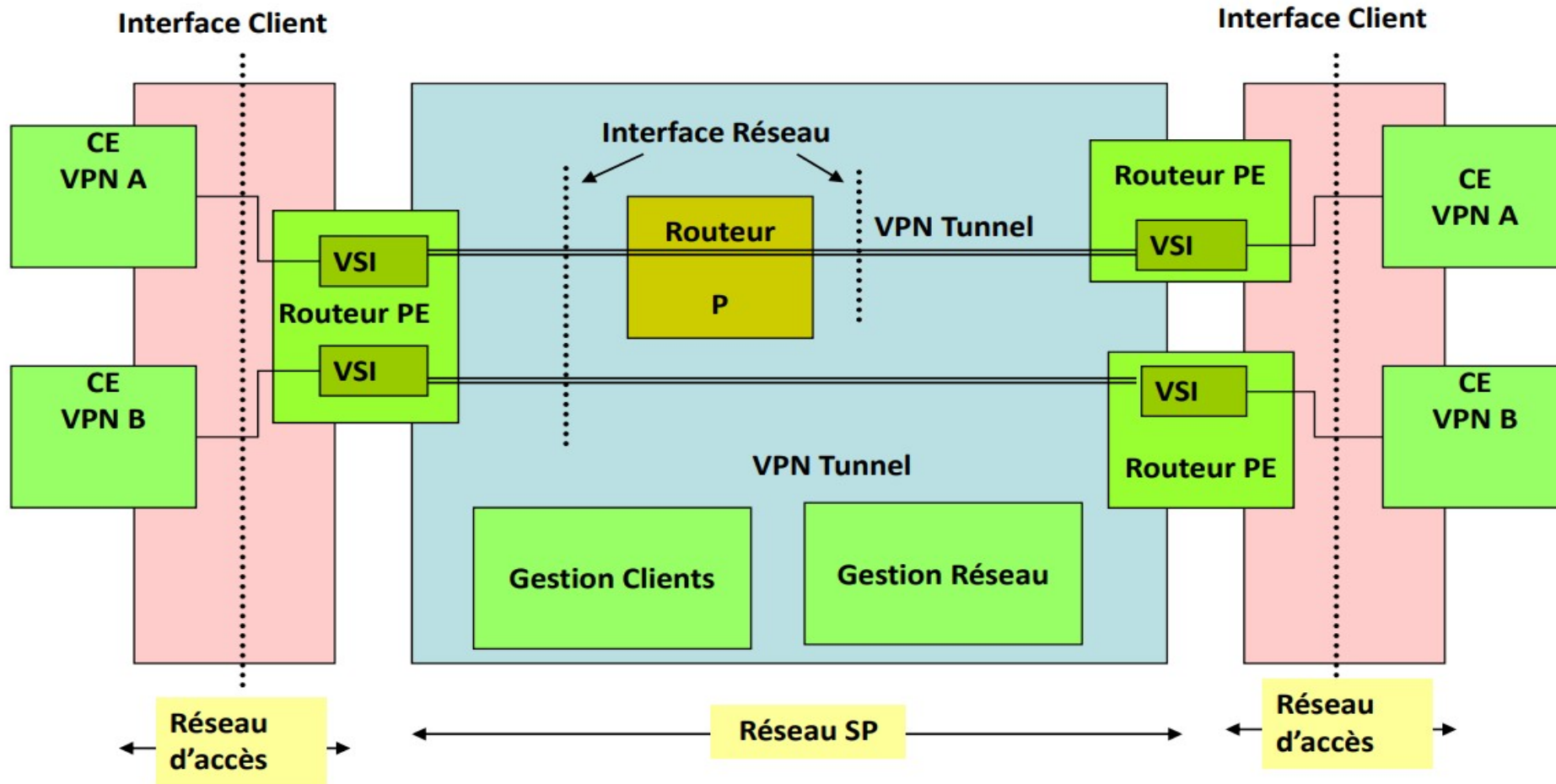


Vis-à-vis des mécanismes de contrôle du VPN le réseau du SP est complètement transparent.

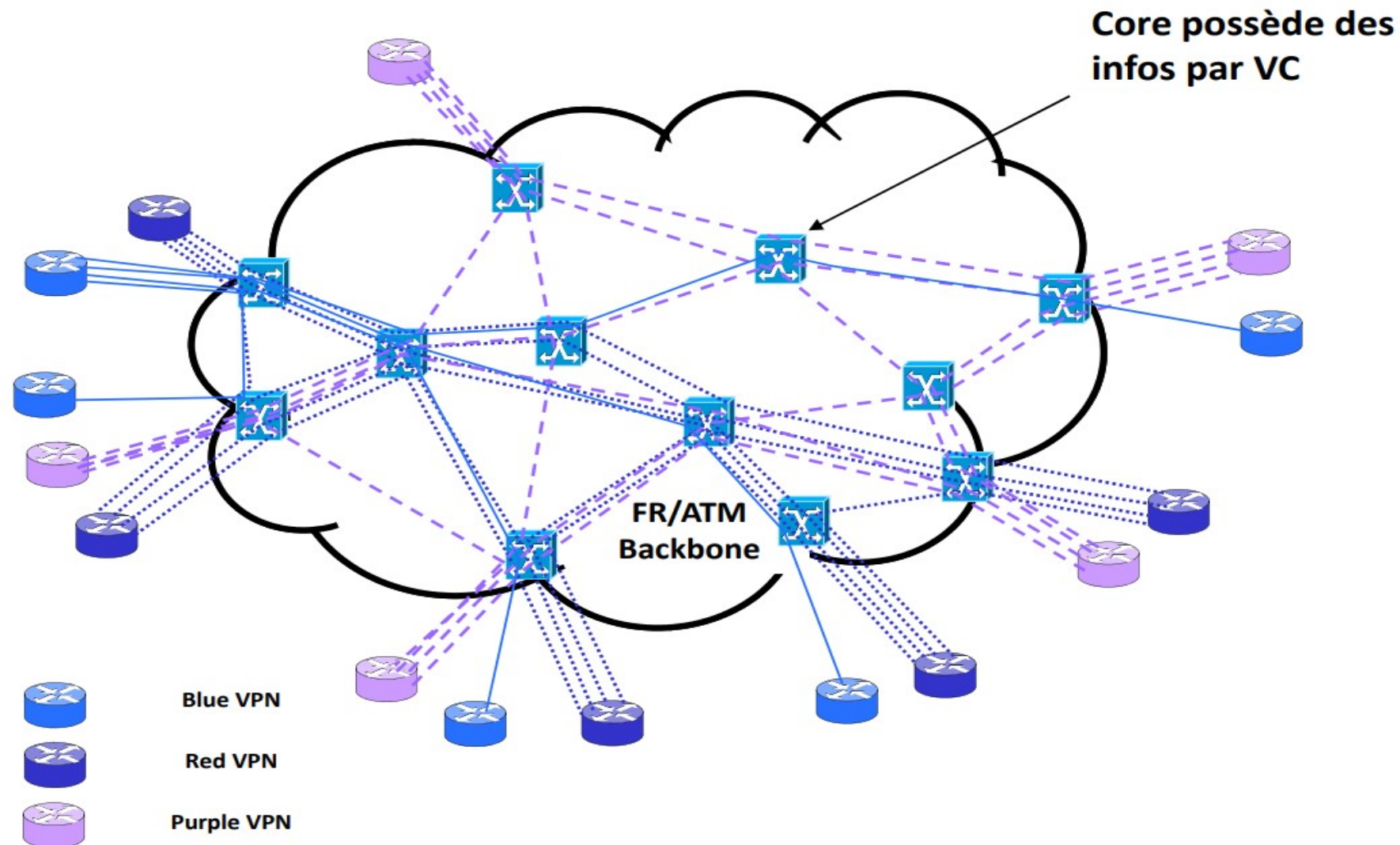
Network-Based layer 2 IPVPN

- ❑ Variante du modèle *Network-Based layer 3*.
- ❑ Le concept de VSI 'Virtual Switching Instance' prend la place de la VFI
- ❑ VSI (PE) supporte les fonctions relatives au
 - ❑ processus de *Forwarding*
 - ❑ des trames de niveau 2,
 - ❑ des cellules ou des paquets pour un VPN spécifiques
 - ❑ en plus de la terminaison des tunnels VPNs.

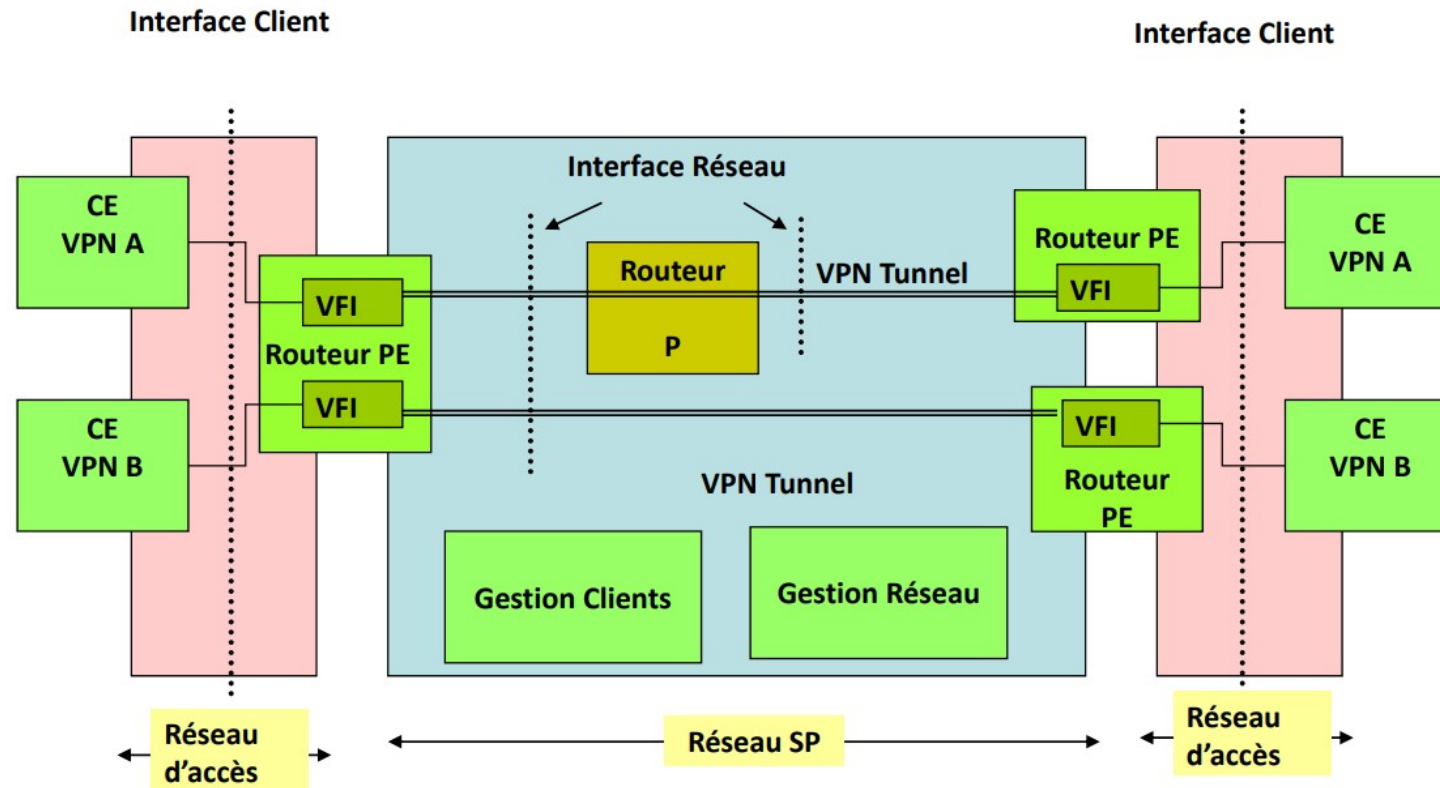
Le modèle Network-Based layer 2 IPVPN



Layer 2 VPNs – FR & ATM



Le modèle Network-Based layer 3 IPVPN



Le modèle Network-Based layer 3 IPVPN

- ❑ Toute l'intelligence est hébergée dans le routeur PE du SP.
- ❑ Chaque routeur PE met en œuvre une ou plusieurs instances VFI et maintient un état par VPN.
- ❑ VFI '*VPN Forwarding Instance*' une entité logique, dédiée, résidant dans le routeur PE qui contient la base d'information du routeur ainsi que la base d'information de transfert (Forwarding) pour un VPN spécifique.
- ❑ VFI termine les tunnels d'interconnexion avec les autres VFIs et peut aussi terminer les connexions d'accès aux CEs.

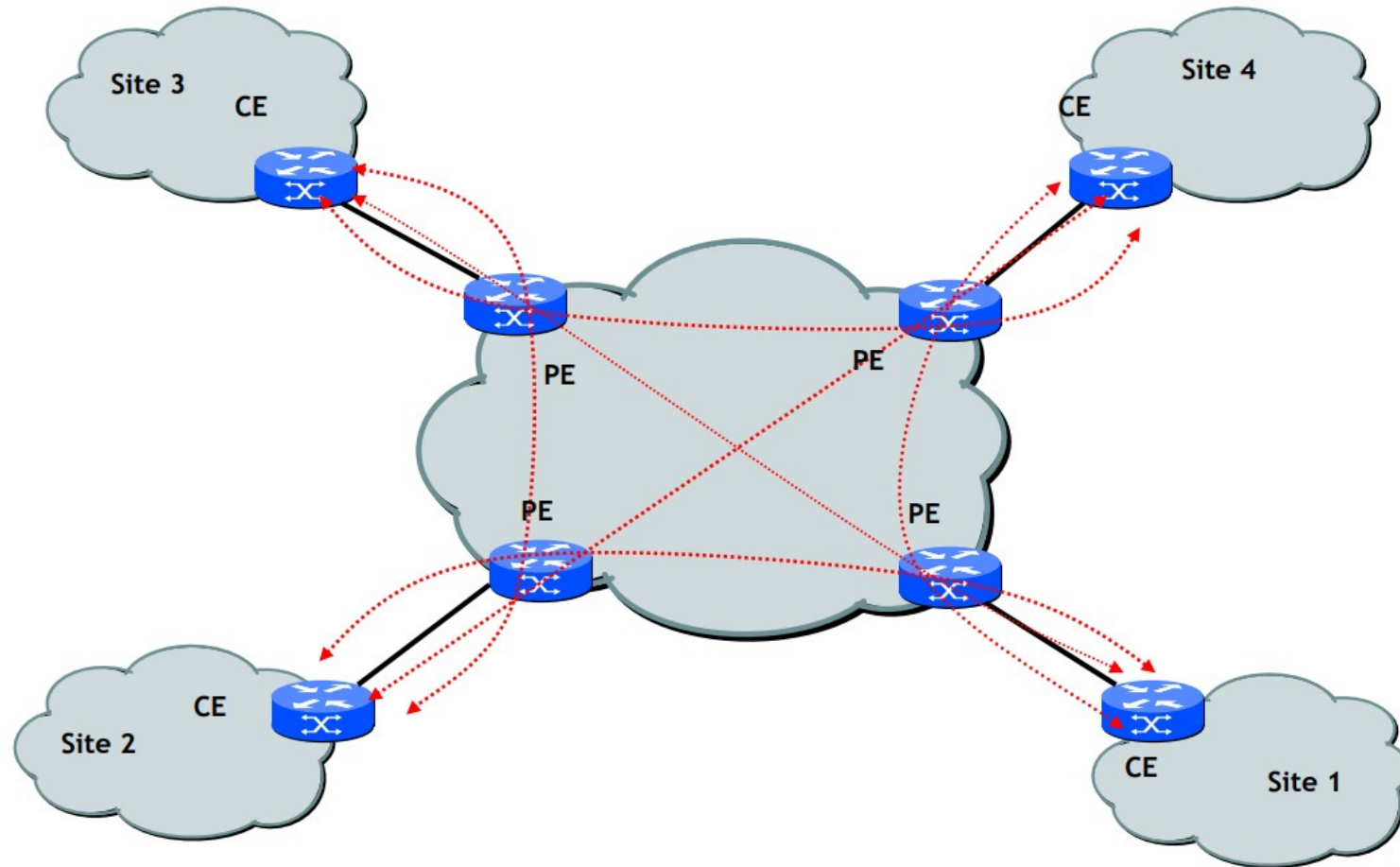
Modèles d'implémentation

- ❑ La logique d'échange des informations de routage entre le CPE et le SP.
- ❑ Deux principaux modèles :
 - ❑ Overlay
 - ❑ Peer to Peer

Le modèle Overlay

- ❑ Le service VPN est fonctionnellement équivalent à des liaisons louées émulées.
- ❑ Le SP et le CE n'échangent aucune information de routage de niveau 3.
- ❑ Offre une nette séparation des domaines de responsabilités Clients et SP.
 - ❑ mise à jour de la matrice de trafic,
 - ❑ l'ajout de (n-1) PVC,
 - ❑ révision de la taille des PVCs en full Mesh,
 - ❑ mise à jour du routage
 - ❑ reconfiguration de chaque CPE pour la topologie couche 3 !

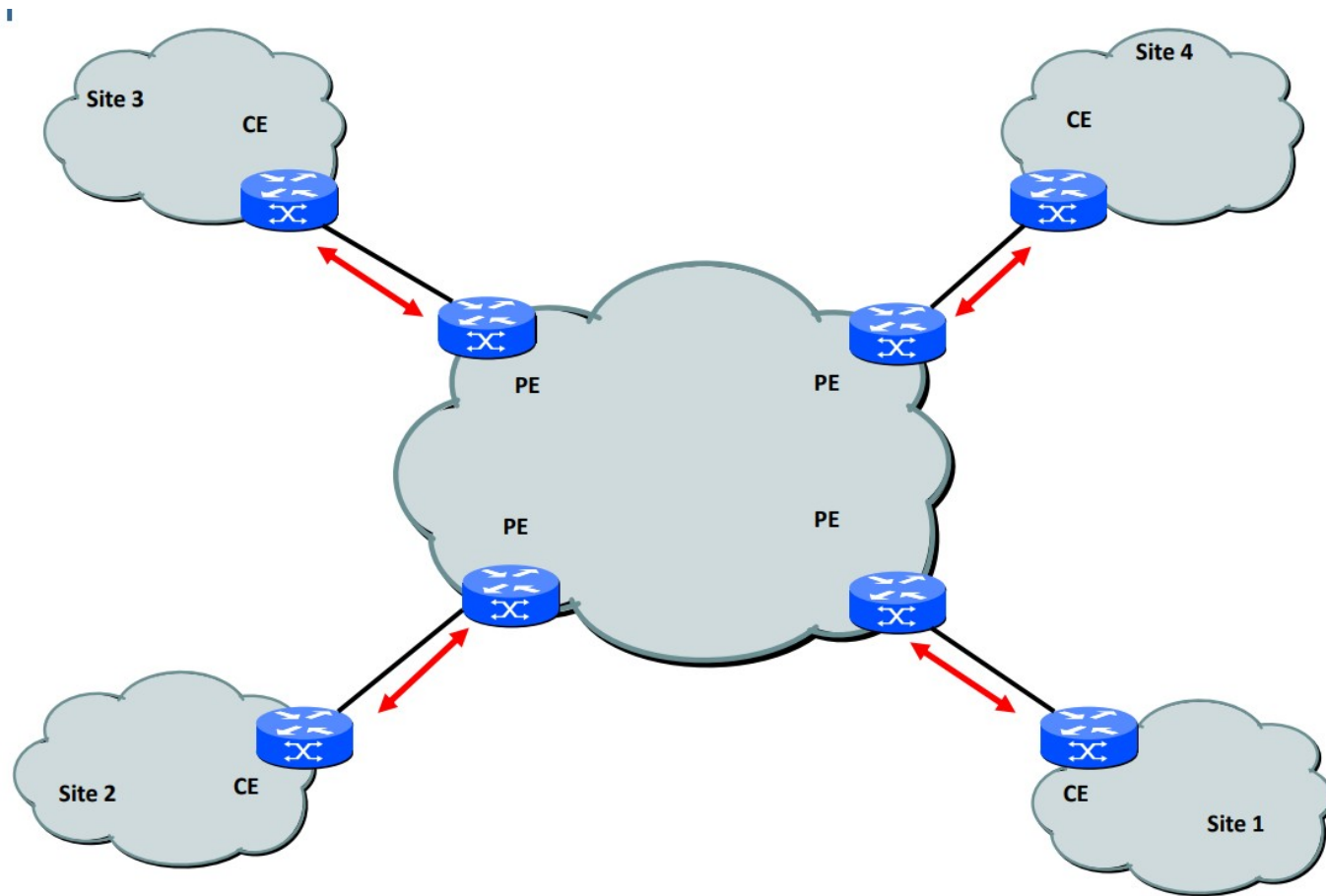
Modèle Overlay VPN



Le modèle Peer to Peer VPN

- ❑ Participation active du SP dans le routage client, l'acceptation de ses routes, leur transport sur le backbone et finalement leur propagation aux autres sites clients.
- ❑ Simple, souple et extensible offrant une facilité de Provisioning des clients et une meilleure gestion des ressources réseau.

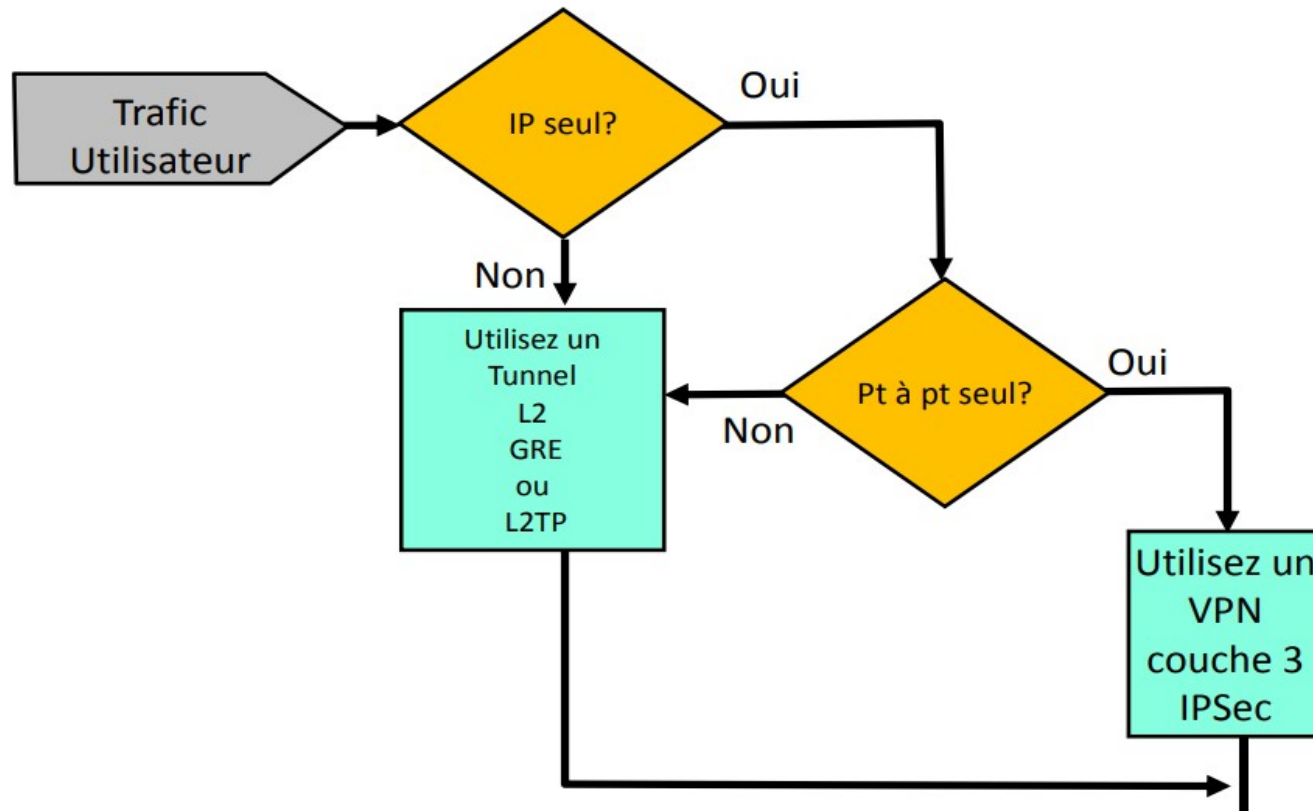
Le Modèle Peer to Peer



VPN (logiciels) et couche OSI

Application	HTTP-s, ...	IKE (ISAKMP, ...)
Présentation		
Session	SSL / TLS	
Transport	TCP	UDP
Réseau	IP / IPSec (ESP, AH)	
Liaison	PPTP, L2TP, ...	
Physique		

Choix de la meilleure technologie VPN



- Sélectionnez la meilleure technologie VPN pour fournir une connectivité réseau selon les besoins du trafic.

Avantages:

☐ **Peer to Peer VPN**

- ☐ Le Fournisseur de service participe au routage client
- ☐ Le réseau client et le réseau du Fournisseur sont bien isolés
- ☐ Un routage optimal entre Les sites clients

☐ **Overlay VPN**

- ☐ Le Fournisseur de service ne participe pas au routage client
- ☐ Configuration usuelle (traditionnelle) pour l'ajout de VNs additionnels
- ☐ Les sites sont les seuls à être provisionnés, non les liaisons
- ☐ Un routage peut devenir non optimal entre Les sites clients

Limitations:

☐ **Overlay VPN**

- ☐ Le routage optimale requiert des VC en full mesh
- ☐ Le provisionning des VCs est manuel
- ☐ La BP doit être provisionnée sur la base du site à site
- ☐ Souffre de l'encapsulation d'overheads

☐ **Peer to Peer VPN**

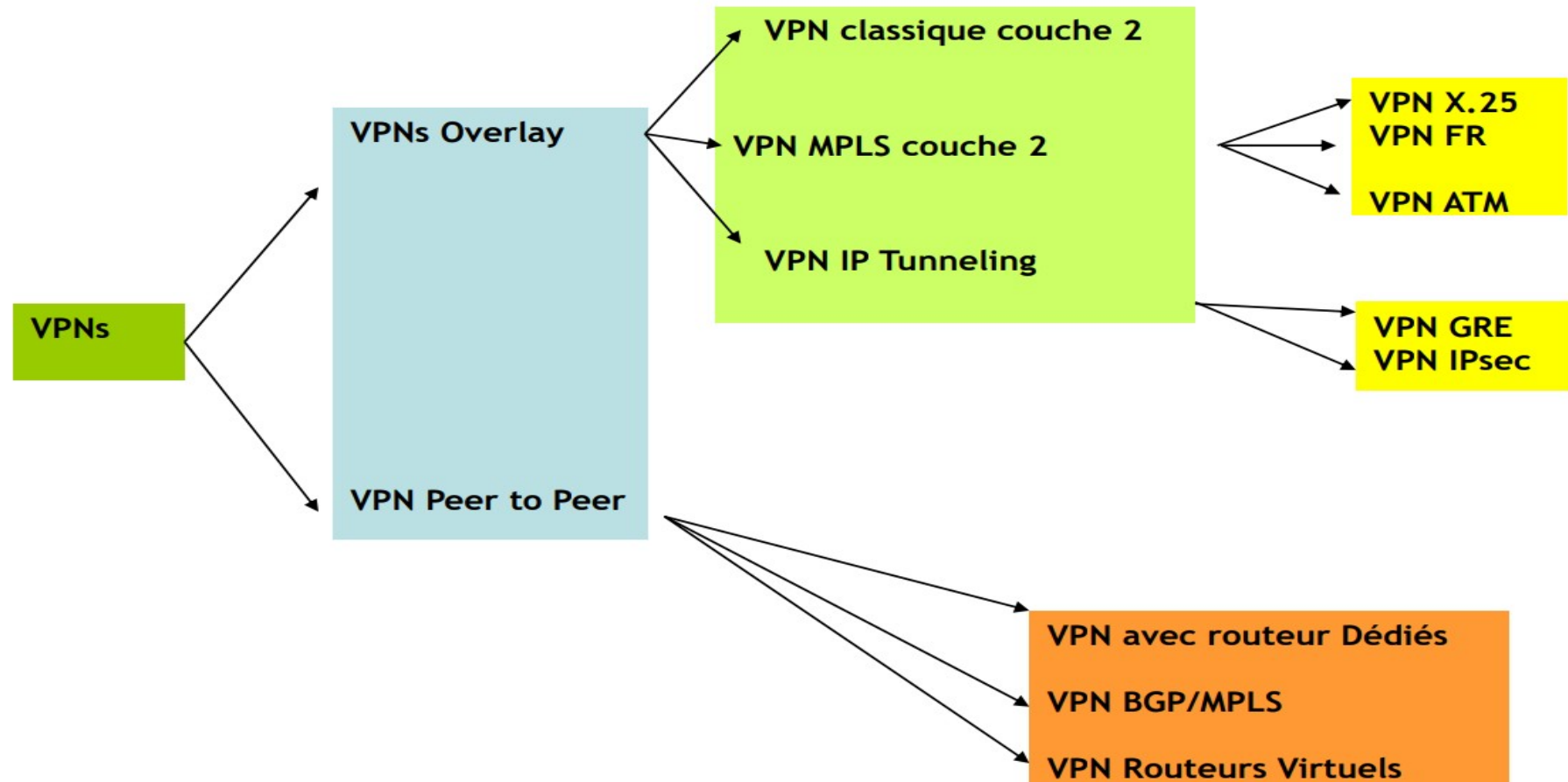
- ☐ Le SP participe dans le routage client
- ☐ Le SP est responsable de la convergence chez le client.
- ☐ Les PEs portent toutes les routes des clients
- ☐ Le SP a besoin d'une connaissance accrue du routage IP

La solution

**Les opérateurs, militent en faveur de
l'utilisation de l'architecture
hybride :**

**Network-Based layer 3 VPN-IP
utilisant une logique Peer to Peer**

Classification VPN-IP :



Classification

- ❑ Classification des VPN en fonction des équipements mis en œuvre (*suite*)
 - ❑ VPN « session » (*Session Based*)
 - ❑ Solutions plus connues sous le nom de VPN SSL
 - ❑ Secteur en plein essor !
 - ❑ Solutions assimilables à des *reverse proxy* HTTPS (i.e. tunnel SSL et redirection de port)
 - ❑ Usage abusif du terme VPN ?
 - ❑ Accès à un portail d' applications et non au réseau interne
 - ❑ Excepté OpenVPN ?
 - ❑ Convient bien aux accès distant

Technologies VPN

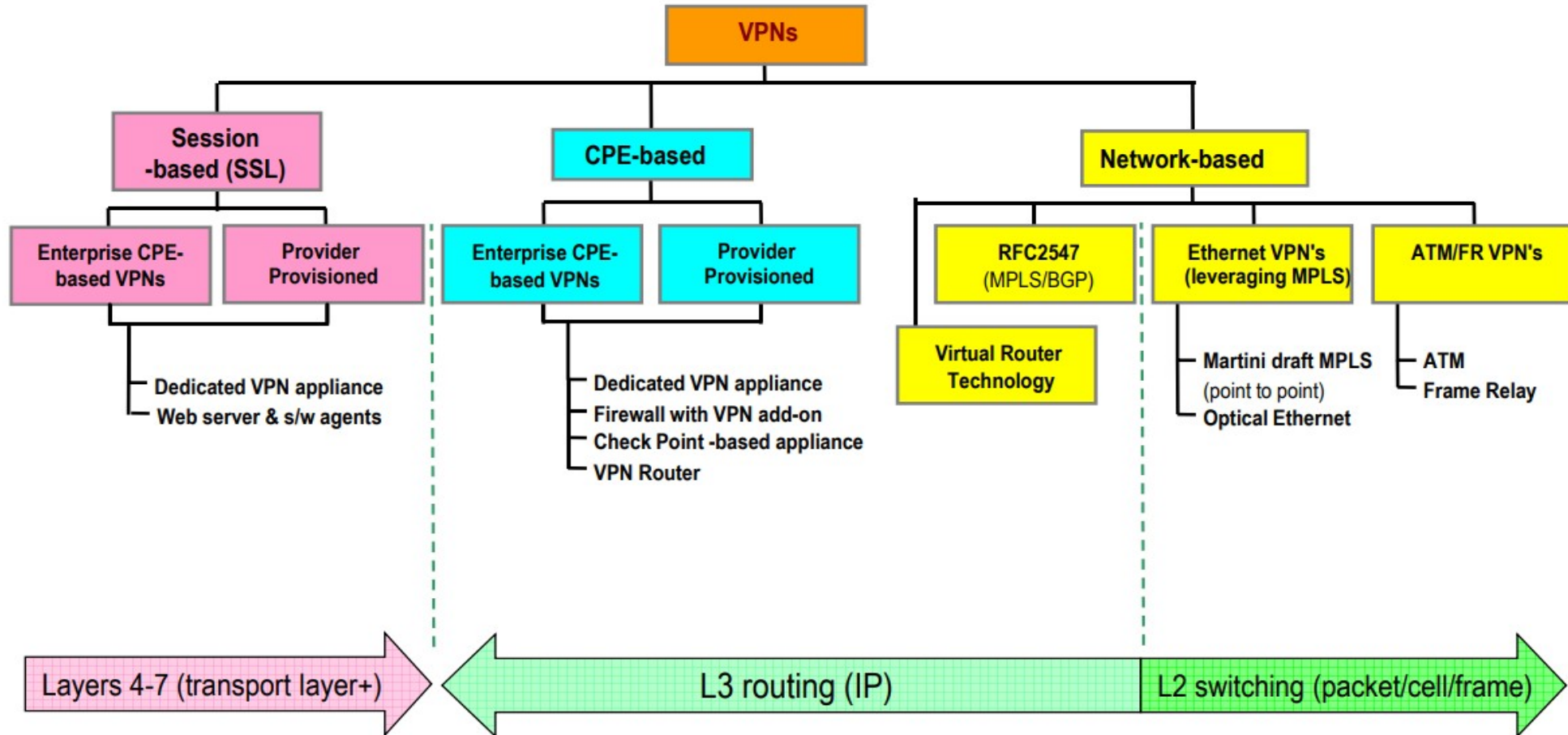
- ❑ VPN

- ❑ Virtual Private Network
 - ❑ Réseau IP privé transporté par un réseau partagé

- ❑ Technologies de VPN

- ❑ VPN sur niveau 1 partagé
 - ❑ VLAN, SDH
 - ❑ VPN sur niveau 2 partagé
 - ❑ Frame Relay, ATM
 - ❑ VPN sur niveau 3 partagé
 - ❑ Tunnels : GRE, L2TP, PPTP, L2F, IPSec
 - ❑ VPN-IP : Tag VPN, MPLS VPN

Synthèse des solutions VPN:



Source Burton Group (d'après Nortel)

VPN : Contraintes des VPN-IP

1. Étanchéité des plans d'adressage et des VPNs

- ☐ Recouvrement possible des adresses utilisées par deux clients : Overlapping
- ☐ Recouvrement possible des adresses clients/backbone
- ☐ Communication any-to-any dans un VPN
 - ☐ Hub & Spoke est toujours possible
- ☐ Backbone vu comme niveau 2 totalement maillé
 - ☐ Les TTL des paquets IP des clients sont préservés
 - ☐ Les routeurs backbone n'apparaissent pas dans traceroute

VPN : Implications techniques

1. Multiples tables de routage

- ❑ Les routeurs de backbone doivent maintenir des tables de routage indépendantes pour chaque VPN
- ❑ => VRF : VPN Route Forwarding table

2. Protocole dynamique d'échange de routes VPN

- ❑ Les routeurs doivent s'échanger les routes de leurs VRF
- ❑ Ce protocole doit savoir différencier les routes par VPN
- ❑ => MP-iBGP

3. Différentiation des paquets IP reçus

- ❑ Les routeurs doivent savoir attribuer les paquets reçus à un VPN
- ❑ => encapsulation des paquets issus de VPNs

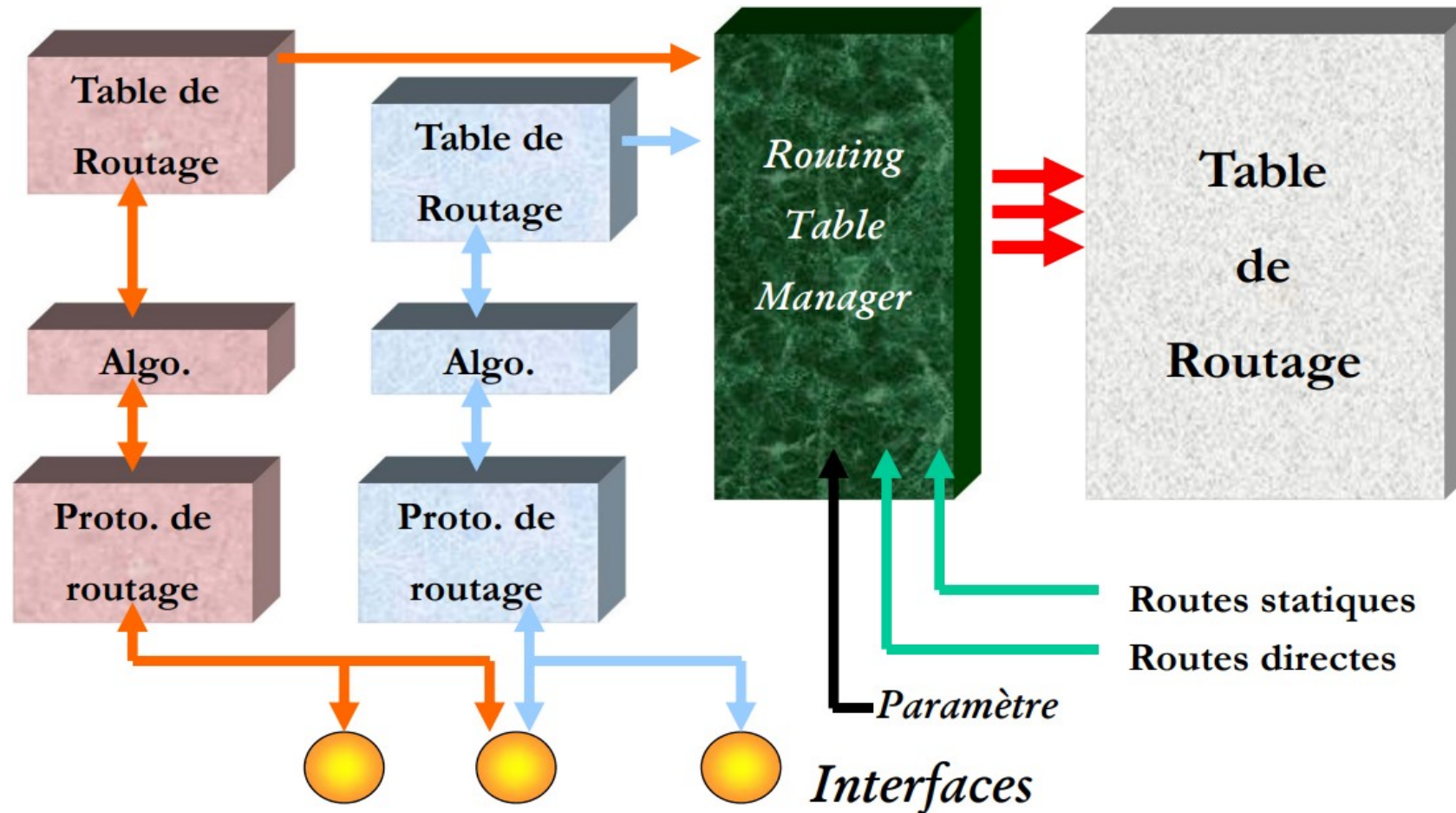
MPLS et VPN-IP

- ❑ Le RFC 2547bis définit un mécanisme permettant aux SPs l'utilisation de leur backbones IP pour offrir des services VPN.
- ❑ Ces VPNs son communément appelés **BGP/MPLS-VPNs** :
 - ❑ Vu que le protocole BGP est utilisé pour la distribution des informations de routage à travers le backbone
 - ❑ Vu que MPLS est utilisé pour acheminer le trafic d'un site VPN à un autre.

Les objectifs

- ❑ Rendre le service simple d'usage pour les clients même s'ils ne disposent pas d'expérience dans le routage IP
- ❑ Rendre le VPN extensible et flexible pour faciliter un déploiement à grande échelle
- ❑ Permettre au SP d'offrir un service à forte valeur ajoutée capable de galvaniser sa loyauté.

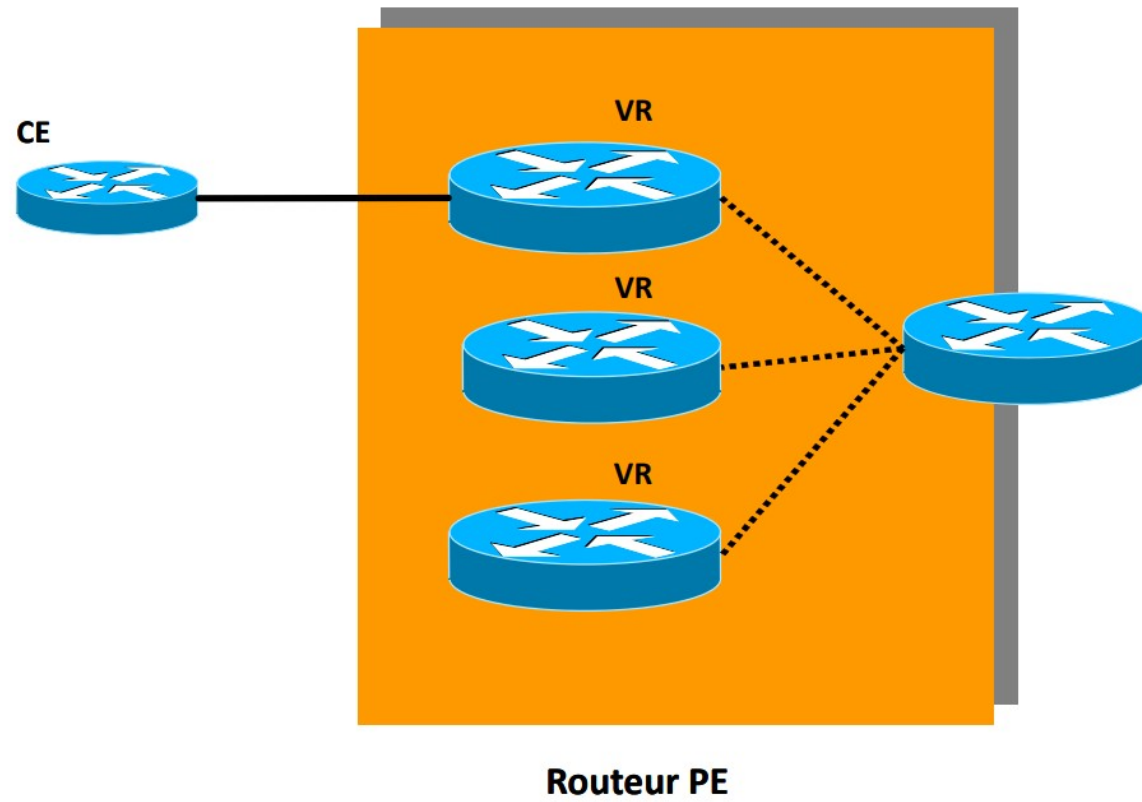
Réseau et table de routage



Composantes du réseau

- ❑ Dans le contexte du RFC 2547, un VPN est défini par une collection de politiques qui contrôlent la connectivité d'un ensemble de sites.
- ❑ Un site client est connecté au réseau du SP par un ou plusieurs ports
- ❑ Le SP associe à chaque port une table de routage VPN appelée ***VRF 'VPN Routing and Forwarding'***.

Virtual Router

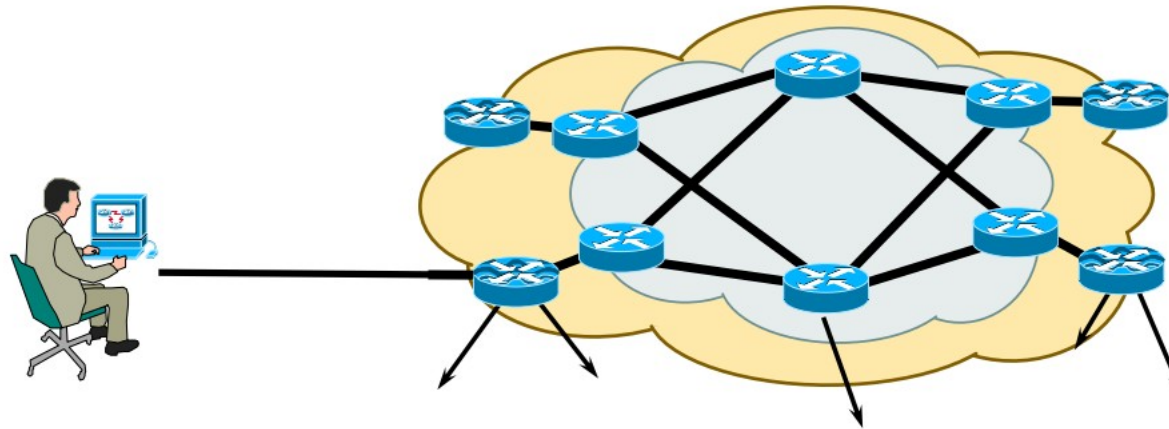


VRF et Multiple Routing Instances

- ❑ VRF : VPN Routing and Forwarding Instance
 - ❑ *VRF Routing Protocol Context*
 - ❑ *VRF Routing Tables*
 - ❑ *VRF CEF Forwarding Tables*

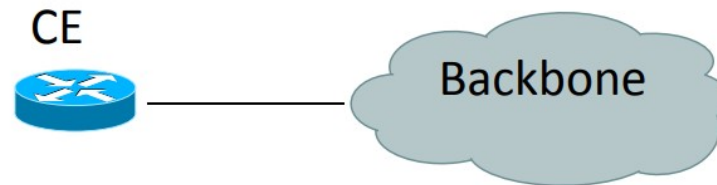
Les composants VPN

- ❑ La mise en œuvre du service VPN repose sur :
 - ❑ Niveau Infrastructure
 - ❑ Niveau gestion.



Customer Edge

- ❑ Le CE (Customer Edge) est l'équipement qui permet l'accès du client au réseau du SP sur une liaison de données à un ou plusieurs routeurs PE.
- ❑ Typiquement, c'est un routeur qui établit une adjacence avec les PEs où il est directement connecté.



- ❑ Après l'établissement de l'adjacence, le CE annonce au PE les routes VPN du site local pour qu'il puisse recevoir de ce dernier les routes VPN distantes.

Provider Edge

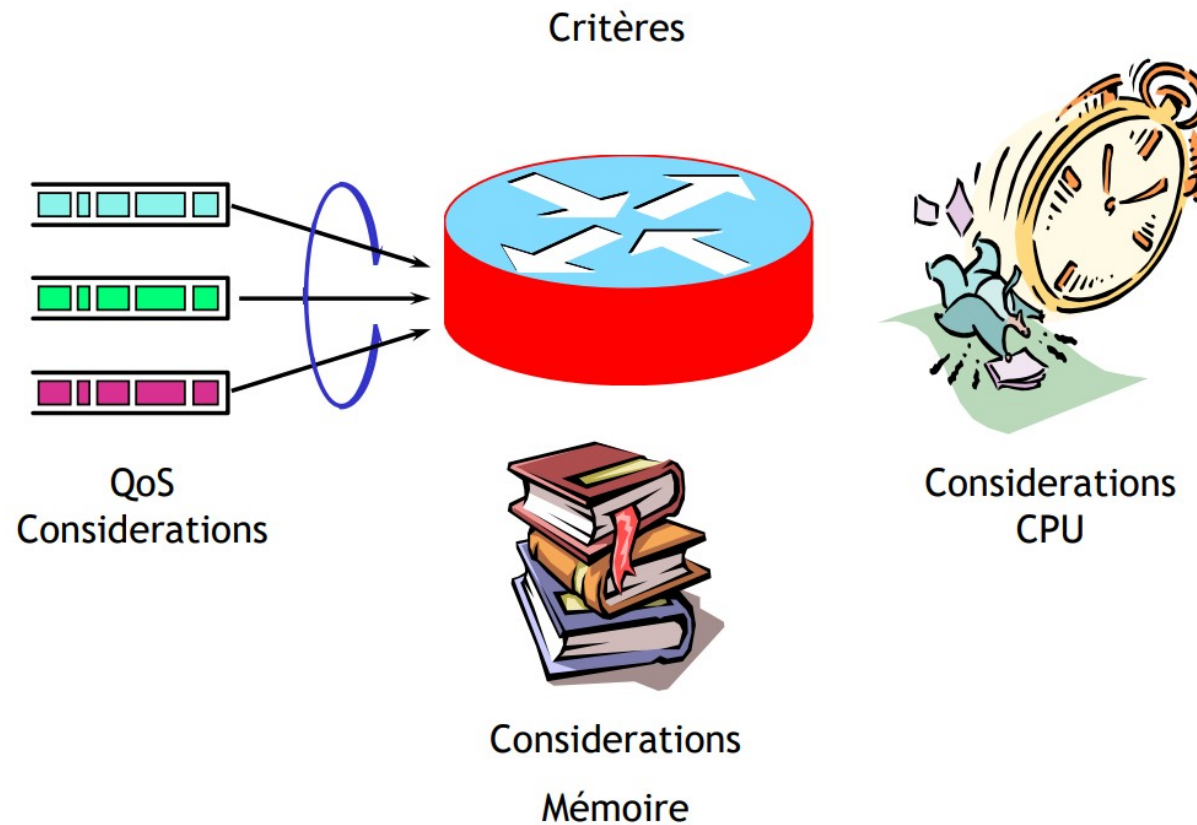
- ❑ C'est grâce au concept de routeur PE (Provider Edge) implémentant une ou plusieurs VRF qu'il est devenu possible de mettre en place des réseaux d'opérateurs souples et commercialement viables.
- ❑ **Le PE échange les informations de routage avec le CE par le routage statique, RIPv2, OSPF ou eBGP.**
- ❑ **Au moment où le PE maintient les informations de routage VPN, il lui est requis seulement de maintenir les routes VPN de ceux qui lui sont directement connectés.**
 - ❑ Ce qui contribue fortement à l'amélioration de l'extensibilité du réseau.

Provider Edge

- ❑ Chaque PE maintient un VRF à chaque site directement connecté.
- ❑ Chaque connexion (Frame Relay, LL, ..) est mappée à une VRF spécifique.
 - ❑ **D'où la justification du choix d'un port (interface), et non un site, pour l'associer à un VRF.**
- ❑ **L'étanchéité des VPNs est rendu possible grâce au maintien par le PE d'une multitude de VRF, rendant la tâche aisée la tâche de routage et améliorant les performances des équipements de commutation.**
- ❑ Les informations de routage locales aux CE sont utilisées par les routeurs PE pour établir, via iBGP, la connectivité IP.
- ❑ **Cette connectivité sera par la suite traduite en connectivité de label.**

Choix du Provider Edge (PE)

■

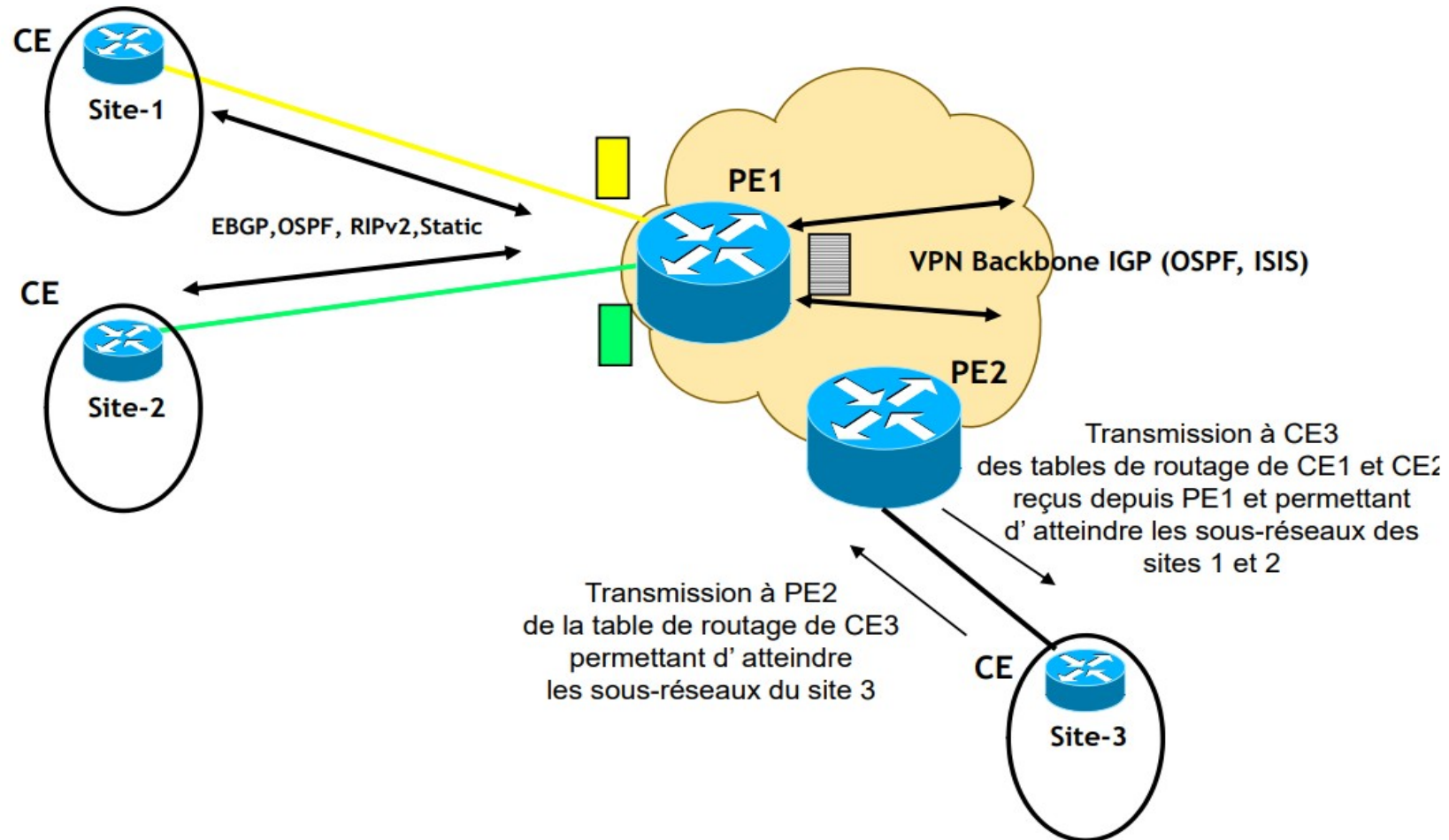


Routage entre CE et PE (1)

□ **Comment distribuer les routes apprises ?**

1. En utilisant un protocole de routage dynamique
 - *Exemple : RIP, OSPF, etc.*
2. Chaque CE d'un VPN(a) doit transmettre à son ou ses PE paires
 - *Les routes apprises pour atteindre les sous-réseaux localisées sur son site*
3. En retour, le ou les PE paires doivent annoncer au CE en question
 - *Les routes à suivre pour atteindre les sous-réseaux localisés sur d'autres sites, mais appartenant au même VPN(a)*
 - *Ces routes seront transmises depuis d'autres PE connectés à des CE du même VPN(a)*

Routage entre PE et CE (2)



Le routeur P

- ❑ Le routeur P n'est connecté à aucun CE
- ❑ Le routeur P est n'importe quel routeur situé dans le backbone, qui n'est connecté à aucun CE.
- ❑ Le P fonctionne en MPLS transit (LSRs) quand il achemine un trafic VPN entre PEs.
- ❑ Les P sont chargés du maintien des routes au PE, et non du maintien d'information spécifique de routage pour les sites clients.

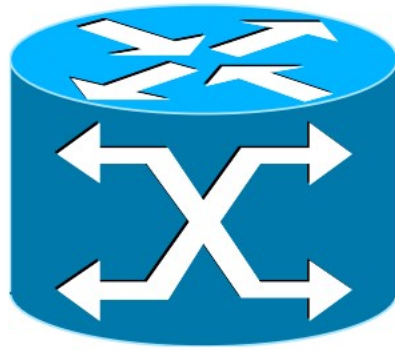
Choix du Provider Router (P)



PE Considerations

Mémoire

Considerations



Considerations

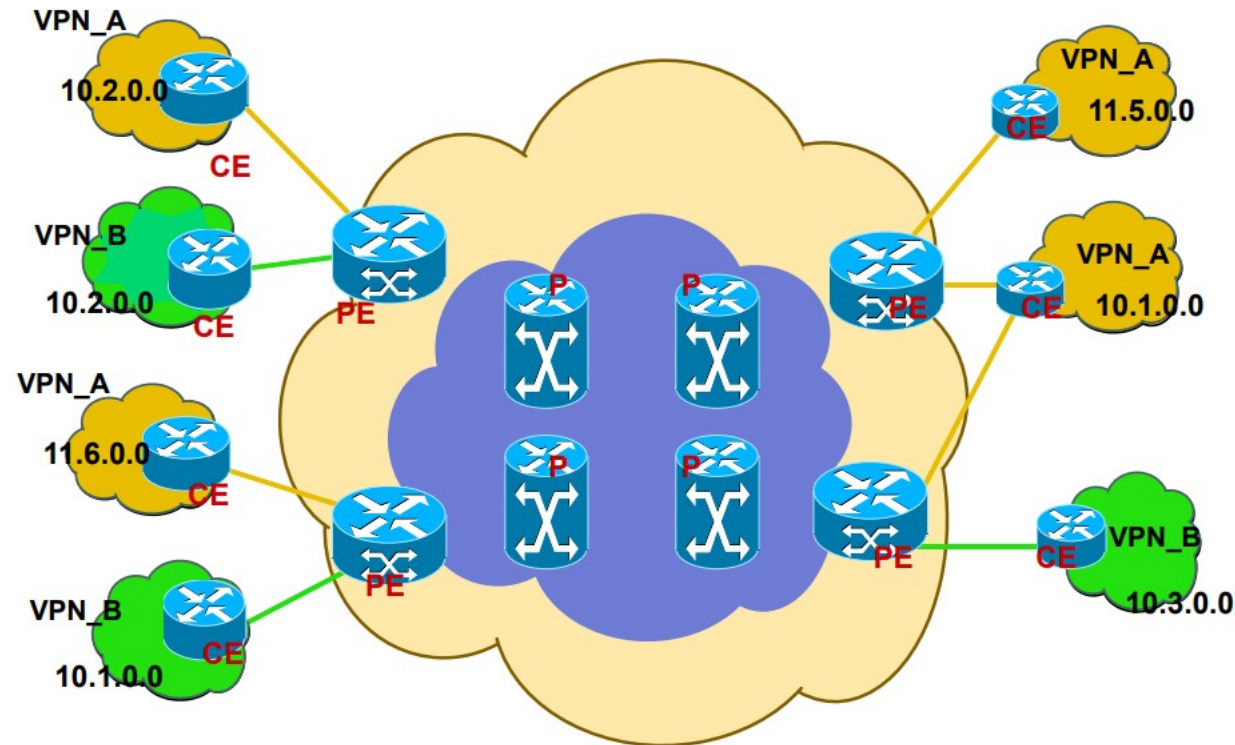
CPU

Routage entre PE (1)

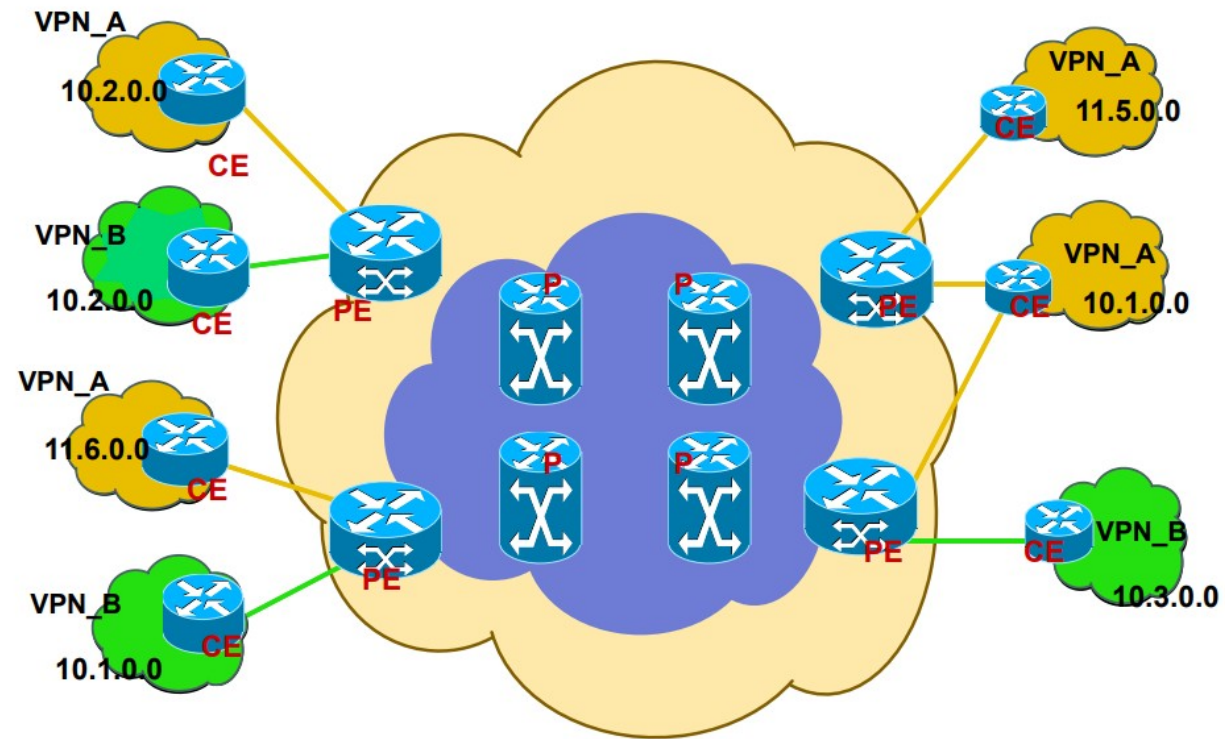
- ❑ **Problème posé**

- ❑ Impossible de distribuer directement les tables de routage des CE

- ❑ *Les sites des différents VPN peuvent en effet utiliser les mêmes systèmes d'adressage privés*



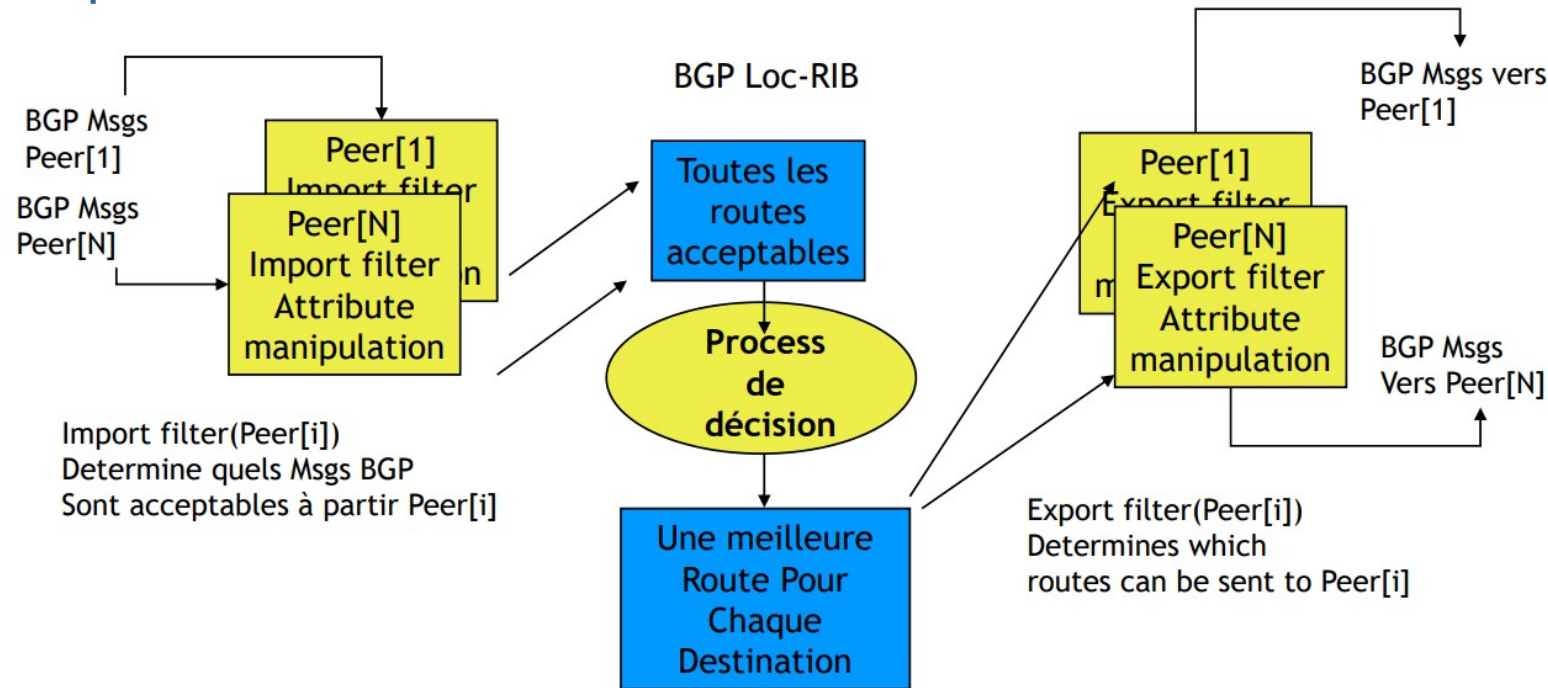
Routage entre PE (2)



BGP Synthèse

- ❑ BGP est un protocole qui peut fonctionner en IGP et EGP
- ❑ Les routeurs internet sont partagés en groupe pour diminuer le trafic de routage
- ❑ Messages BGP
 - ❑ Initialisation
 - ❑ Mise à jour
 - ❑ Notification d'erreur
- ❑ BGP utilise une connexion TCP
- ❑ BGP est riche en attributs : mise en œuvre du routage par politique

Le modèle conceptuel de BGP



BGP Routing Information Base

Contient toutes les routes acceptables apprises à partir des voisins + routes internes

BGP decision process selects

La meilleure route vers chaque destination

Connectivité

- ❑ Le modèle MPLS-VPN se base sur un contrôle plus granulaire des informations de routage par l'ajout de deux mécanismes supplémentaires à savoir :
 - ❑ **le Multiples Forwarding Tables**
 - ❑ Au niveau d'un PE, chacune de ses VRFs est associée à un ou plusieurs de ses ports (interfaces/sous interfaces) qui le connectent directement au site client.
 - ❑ Si un site donné contient des machines qui sont membres dans plusieurs VPNs, la VRF associée au site client contient alors des routes pour tous les VPNs dans lesquels ce site est membre.
 - ❑ **le BGP extended communities.**
 - ❑ La distribution des informations de routage est conditionnée par l'usage des nouveaux attributs du BGP que sont les Extended Communities.

BGP Attributes

- BGP attributes :
 - AS-path *
 - Next-hop *
 - Local preference
 - Multi-exit-discriminator (MED)
 - Origin *
 - Community

*** = Well-known mandatory attribute**

L'algorithme de Sélection de route BGP

Pour une route synchronisée ayant un next-hop valide (pas de boucles) :

1. Utilise **weight** le plus élevé (local au routeur)
2. Utilise **local preference** le plus élevé
3. Le routeur est à l'origine de la route (**network, redistribute puis aggregate**)
4. Utilise le plus court **AS-path** (sauf si BGP 'best path as-path ignore')
5. Utilise le plus faible **Origin code** (IGP < EGP < incomplete)
6. Utilise le plus faible **MED** (pour les as-path qui commence par le meme AS sauf si BGP 'always-compare-med')
7. Utilise **chemin EBGP** puis **chemin IBGP** (interne ou EBGP Intra-confederation)
8. Utilise le chemin utilisant le plus **proche IGP neighbor**
9. Pour l'IBGP le chemin le plus vieux (sauf si BGP 'bestpath compare-routerid'. Il y a load-sharing si maximum-path n')
10. Pour l'IBGP le metric igp le plus faible
11. Utilise le voisin avec le 'BGP router id' le plus faible

Connectivité

- ❑ BGP extended communities

- ❑ Ces attributs font parties des messages BGP en tant qu'attributs de la route.
 - ❑ Ils identifient la route comme appartenant à une collection spécifique de routes et qui font objet de la même politique de traitement.
 - ❑ Chaque attribut BGP Extended Community doit être globalement unique et ne peut alors être utilisé que par un seul VPN.
- ❑ Les attributs BGP Extended Communities utilisés sont de 32 bits.
 - ❑ L'utilisation de ces 32 bits vise l'amélioration de l'extensibilité des réseaux d'opérateurs à 232 communities. Par ailleurs et puisque l'attribut contient l'identificateur du système autonome du SP, il peut aussi servir pour contrôler l'attribution locale tout en maintenant son unicité.

Connectivité

- ❑ Trois types d'attributs BGP Extended Communities sont utilisés :
 - ❑ Le **RT (Route Target)** qui identifie une collection de sites VFRs à qui le PE distribue les routes. Un PE utilise cet attribut pour contraindre l'import de routes vers ses VRFs.
 - ❑ Le **VPN-of-origin** qui identifie une collection de sites et établit la route associée comme venant d'un des sites de l'ensemble.
 - ❑ Le **Site-of-origin** qui identifie le site spécifique duquel le PE apprend une route. Il est encodé comme attribut de « **route origin extended community** », qui peut être utilisé pour prévenir les boucles de routage.

Connectivité

- ❑ En mode opérationnel, et avant de distribuer ses routes locales aux autres PEs, le PE d'entrée affecte un RT à chaque route apprise par les sites directement connectés, qui est basé sur la valeur de la politique cible d'export configurée.
 - ❑ Cette approche offre une flexibilité énorme dans la mesure où un PE peut attribuer un RT à une route.
- ❑ Le PE d'entrée (ingress) peut ainsi être configuré pour assigner un seul RT à l'ensemble des routes apprises d'un site donné, ou d'assigner un RT à un groupe de routes apprises d'un site et autres RT aux autres routes apprises d'un autre.

Connectivité

- ❑ Avant d'installer les routes distantes distribuées par un PE, chaque VRF dans un PE sortant (egress) est configurée avec une politique import cible.
 - ❑ Un PE peut uniquement installer une route VPNv4 dans une VRF si le RT transporté avec la route correspond à une VRF import cible.
- ❑ Cette approche permet à un SP d'utiliser un seul mécanisme pour servir les clients ayant une large politique de connectivité inter sites. Par le biais d'une configuration sereine d'Import Target et Export Target, le SP peut alors construire différents types de topologies VPN :
 - ❑ maille complète,
 - ❑ maille partielle,
 - ❑ Hub,
 - ❑ Spoke,

Routage entre PE

- ❑ **Solution proposée**

- ❑ Utiliser des extensions de BGP (Border Gateway Protocol)
 - ❑ *MP-iBGP = entre PE d'un même AS (MP=MultiProtocol, i=interior)*
 - ❑ *MP-eBGP = entre PE de différents AS (e=exterior)*
- ❑ Distribuer non pas des adresses IP, mais des **adresses IP-VPN** <composées de deux éléments
 - ❑ *Un identifiant appelé RD (Route Distinguisher) composé*
 - ❑ D' un identifiant de FAI (ex : numéro AS de système autonome)
 - ❑ D' un identifiant de VPN : RT
 - ❑ *Une adresse de sous-réseau*
 - ❑ Située à l'intérieur du VPN identifié au sein du RD

VPN : MP-iBGP

- ❑ Multiple Protocol - internal BGP
- ❑ Extension du protocole BGP v4
 - ❑ nouvelle capability : vpn_ipv4
 - ❑ AFI/SubAFI = 1/128 identifie la famille de routes vpn_ipv4
 - ❑ Rappel : AFI/SubAFI =1/1 identifie la famille de routes ipv4 unicast
- ❑ Nouveau format de route
 - ❑ Route Distinguisher + préfixe IP v4
- ❑ Fonctionnement similaire à l'iBGP classique
 - ❑ IGP nécessaire
 - ❑ routeurs iBGP fully meshed
 - ❑ route reflectors ...

Over lapping

- ❑ Ainsi, la famille d'adresses VPNv4 a été adoptée comme solution à la problématique d'overlapping.
 - ❑ Une adresse VPNv4 est formée de 12 Octets.
-
- ❑ 8 octets composent le RD 'Route Distinguisher' suivi des octets de l'adresse IPv4.



Format de l'extension d'adresse VPN v4

VPN : MP-iBGP

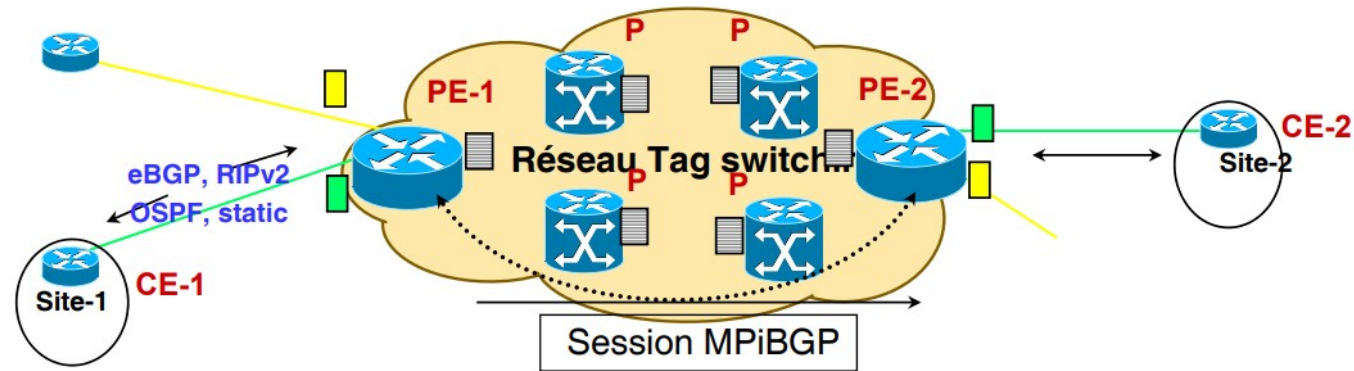
- ❑ Route vpn_ipv4

- ❑ 96 bits (64 + 32)
- ❑ 64 bits pour le Route Distinguisher RD
- ❑ 32 bits pour le préfixe ip_v4

- ❑ Update MP-iBGP

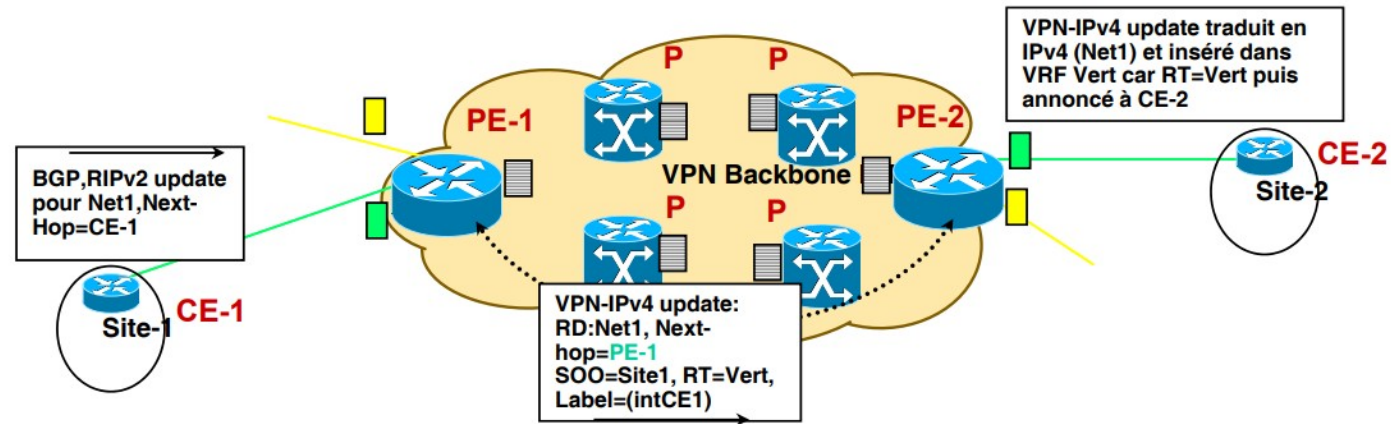
- ❑ route vpn_ipv4
- ❑ les attributs BGP standards (AS_path, Local Pref, MED, ...)
- ❑ attribut SOO Site Of Origin
- ❑ attribut Route Target RT : identifie le VPN
- ❑ attribut Route Origin : identifie le CE origine de la route
- ❑ Tag Externe : émit par le PE d'origine

label VPN : MP-iBGP



- ❑ PE-1 apprend les routes du site-1
 - ❑ Elles sont intégrées dans la VRF verte
 - ❑ Celle-ci est identifiée par le Route Target RT
- ❑ La VRF verte est redistribuée dans le domaine MP-iBGP
 - ❑ en ajoutant le RD pour fabriquer une route vpn_ipv4 unique
 - ❑ en ajoutant le RT et le S00 (site of origin)

Label VPN : MP-iBGP



- ❑ PE-2 reçoit la route de PE-1
 - ❑ récupère le préfixe ipv4
 - ❑ l'intègre dans la vrf verte (grâce au RT)
 - ❑ next-hop = PE-1
- ❑ PE-2 redistribue les routes de CE-1 à CE-2
 - ❑ RIPv2, BGP, OSPF
 - ❑ next-hop PE-2

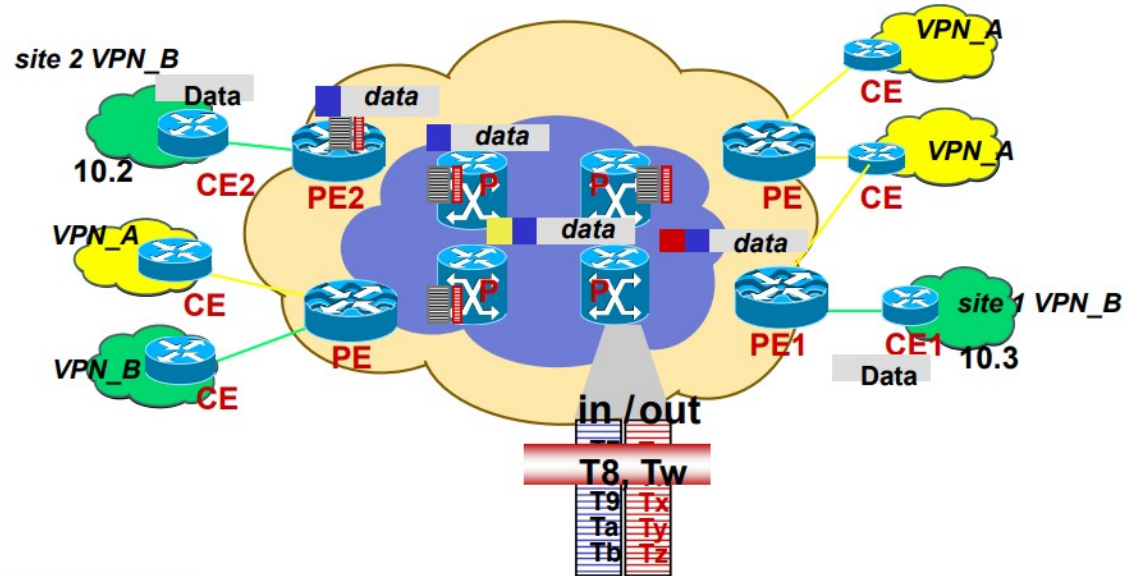
Problématique

- ❑ Comment faire acheminer un paquet issu d'un site VPN vers un autre site du même VPN ?
- ❑ comment partager l'infrastructure public entre des sites privés ?

VPN : Différents label

- ❑ Identification du PE de sortie
 - ❑ Routage inter-PEs
 - ❑ un premier label
 - ❑ Traité par les Ps
 - ❑ Service de routage
- ❑ Identification du routeur virtuel de sortie
 - ❑ Routage inter-VPN
 - ❑ Un deuxième label
 - ❑ Traité par les PEs
 - ❑ Service VPN

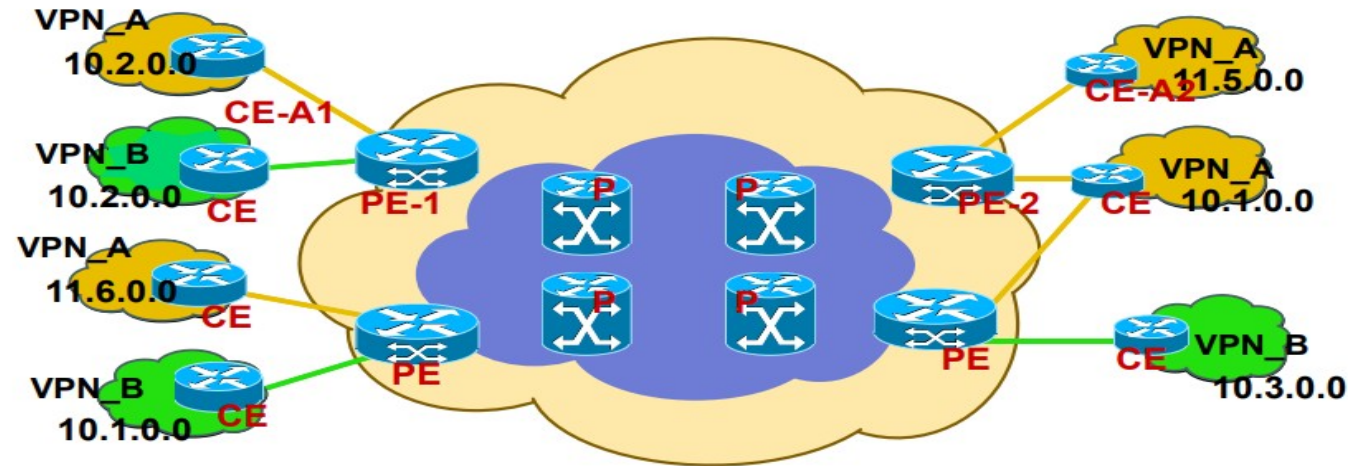
VPN : Différents label



- ❑ **label Interne**
 - ❑ utilisé par Tag Switching (routeurs P et PE)
- ❑ **label Externe**
 - ❑ utilisé pour identifier le VPN d'une trame IP (routeurs PE)
 - ❑ égal au RT

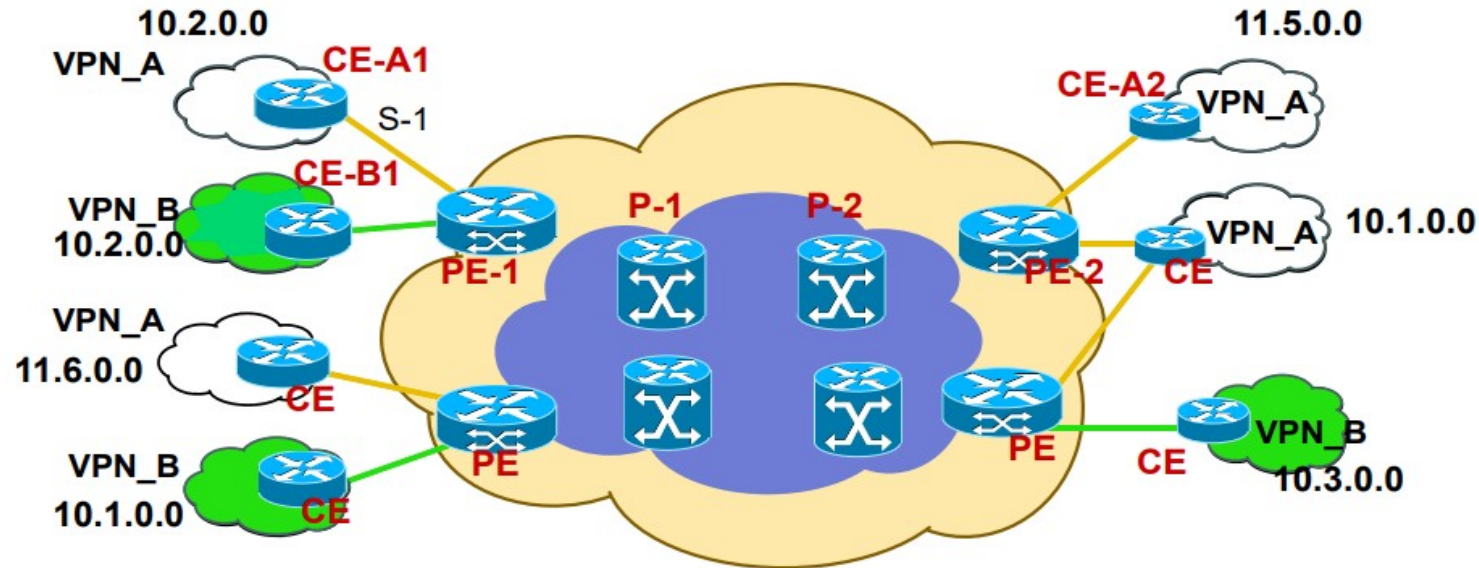
VPN : Envoi d'une trame

T



- ❑ Paquet : 10.2.1.1 vers 11.5.1.1 (VPN-A)
 - ❑ Routé par CE-A1 vers le PE-1
- ❑ Traversée de PE-1
 - ❑ PE-1 connaît le VPN par interface d'entrée => le label externe RT
 - ❑ PE-1 connaît le next-hop PE-2 grâce à la VRF => le label interne (LFIB)
 - ❑ Encapsulation du paquet avec les deux labels
 - ❑ Envoi vers PE-2 à travers le réseau label Switching

Tag VPN : Réception d'une trame



- ❑ Paquet reçu par PE-2
 - ❑ Suppression du label Interne par le P de sortie
 - ❑ Suppression du label Externe => RD récupéré
 - ❑ Avec RD, utilisation du bon VRF
 - ❑ PE-2 route le paquet vers CE-A2

Commutation de labels

- ❑ Utilisation de deux niveaux de labels

- ❑ Un premier niveau de label

- ❑ *Utilisé par les PE pour atteindre les PE apparaissant comme prochain saut dans leurs tables de routage*

- ❑ Un deuxième niveau de label

- ❑ *Utilisé par les PE pour atteindre le bon CE*

- ❑ Chaque PE allouera un label à chaque CE/VPN auquel il est connecté

- ❑ Distribution des labels

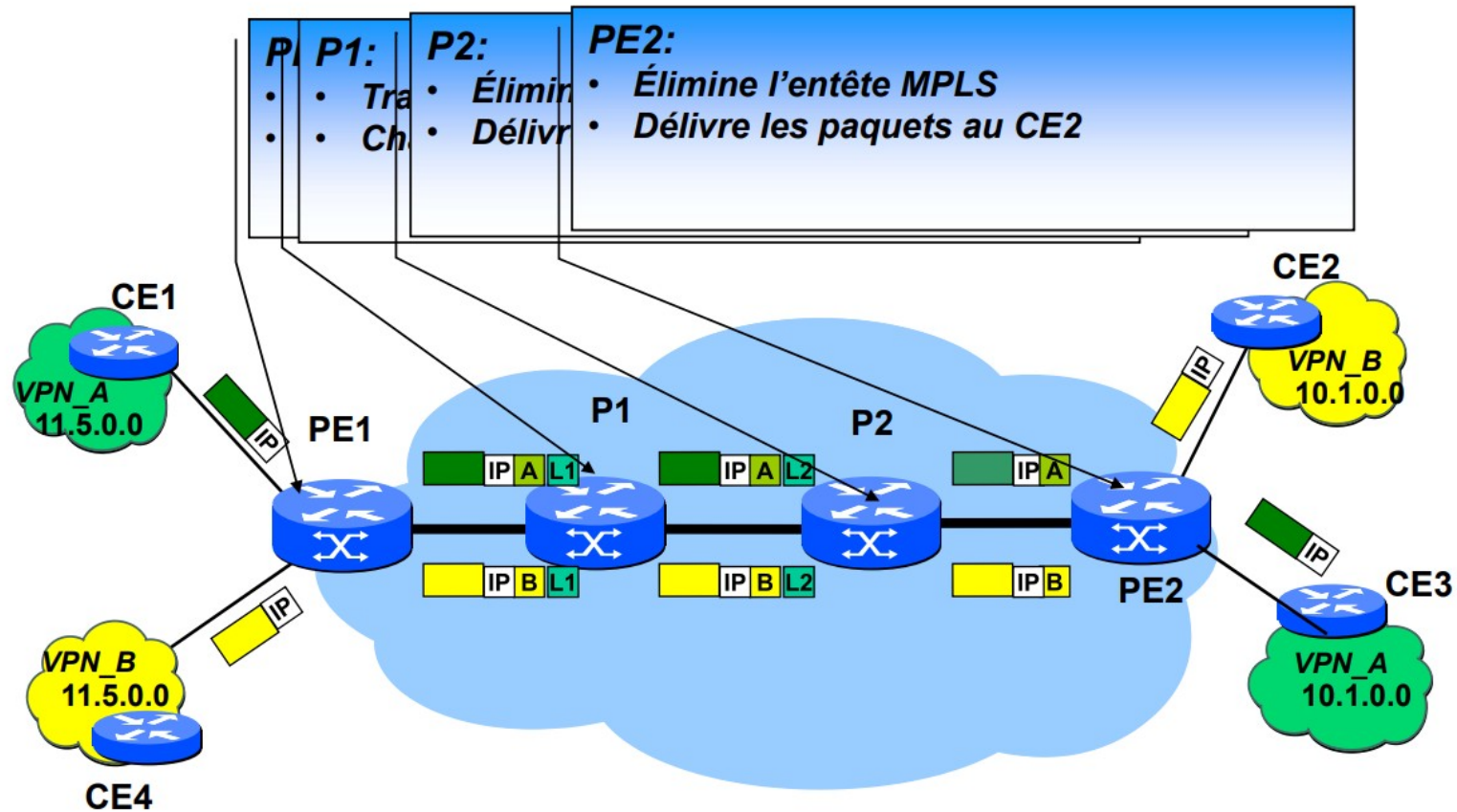
- ❑ Utilisation de LDP à l' intérieur du backbone

- ❑ *Entre routeurs P et PE, ou P et P*

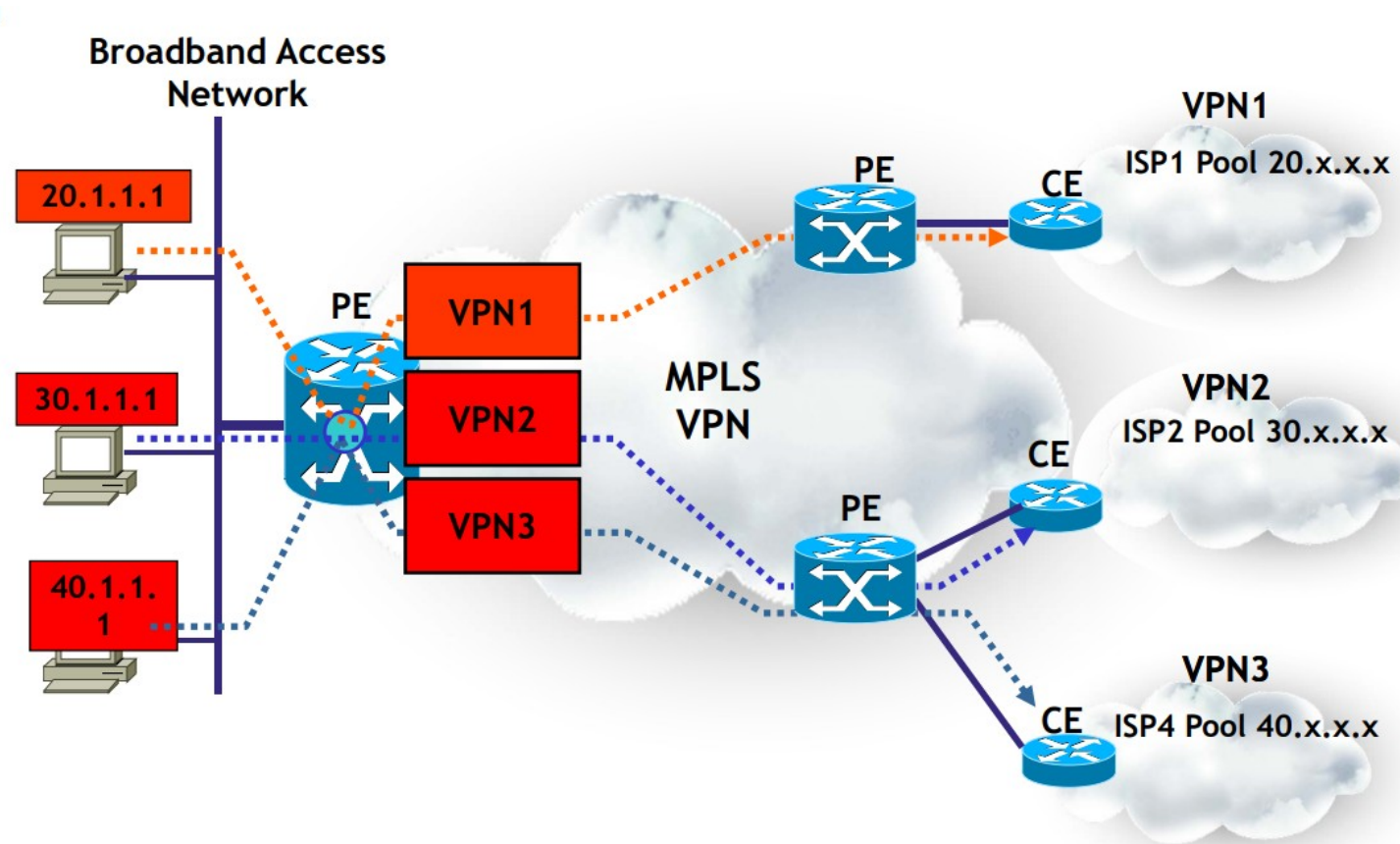
- ❑ Utilisation de MP-BGP

- ❑ *Entre routeurs PE en transmettant simultanément routes et labels*

Architecture MPLS/VPN-IP



VPN Sélection



Apports du MPLS VPN

- ❑ La mise en place du service MPLS-VPN profite pleinement des avantages de l'architecture MPLS en terme séparation des plans logiques et physiques à savoir :
- ❑ L'absence de contraintes sur le plan d'adressage adopté par chaque VPN client. Le client peut ainsi utiliser un adressage publique ou privé. Pour le SP, plusieurs clients peuvent utiliser les mêmes plages d'adresses.
- ❑ Du fait que le modèle adopté est le Network Based layer 3 VPN, le CE n'échange pas directement les informations de routage avec les autres CE du même VPN. Les clients ne sont alors pas obligés d'avoir une parfaite connaissance du schéma de routage utilisé dans le réseau.
- ❑ Les clients n'ont pas un backbone virtuel à administrer. De ce fait, la tâche de management PE ou P, voire même CE, est une tâche de moins.

Apports du MPLS VPN

- ❑ Le SP administre un seul réseau mutualisé pour l'ensemble de ses clients.
- ❑ Le VPN client peut s'appuyer sur un ou plusieurs backbone SP.
- ❑ Même sans l'utilisation de technique de cryptographie, la sécurité est équivalente à celle supportée par les VPNs de la couche 2 (ATM ou Frame Relay).
- ❑ Les SP peuvent utiliser une infrastructure commune pour offrir les services de connectivité VPN et Internet.
- ❑ Conformité au modèle DiffServ.
- ❑ Une QoS flexible et extensible grâce à l'utilisation des bits EXP de l'entête MPLS ou par l'utilisation du trafic Engineering (RSVP)
- ❑ Le modèle RFC 2547bis est indépendant du lien de la couche 2.

Apports du MPLS VPN

- ❑ Absence de contraintes sur le plan d'adressage
- ❑ CE n'échange les informations de routage qu'avec le PE
- ❑ Le management pour les clients, est une tâche de moins
- ❑ SP administre un seul réseau mutualisé
- ❑ Le VPN client peut s'appuyer sur un ou plusieurs backbones SP.
- ❑ Sécurité est équivalente à celle supportée par les VPNs de la couche 2
- ❑ Même infrastructure pour les services de connectivité VPN et Internet.
- ❑ Support de la QoS
- ❑ Le modèle RFC 2547bis est indépendant du lien de la couche 2.

Les principales tendances

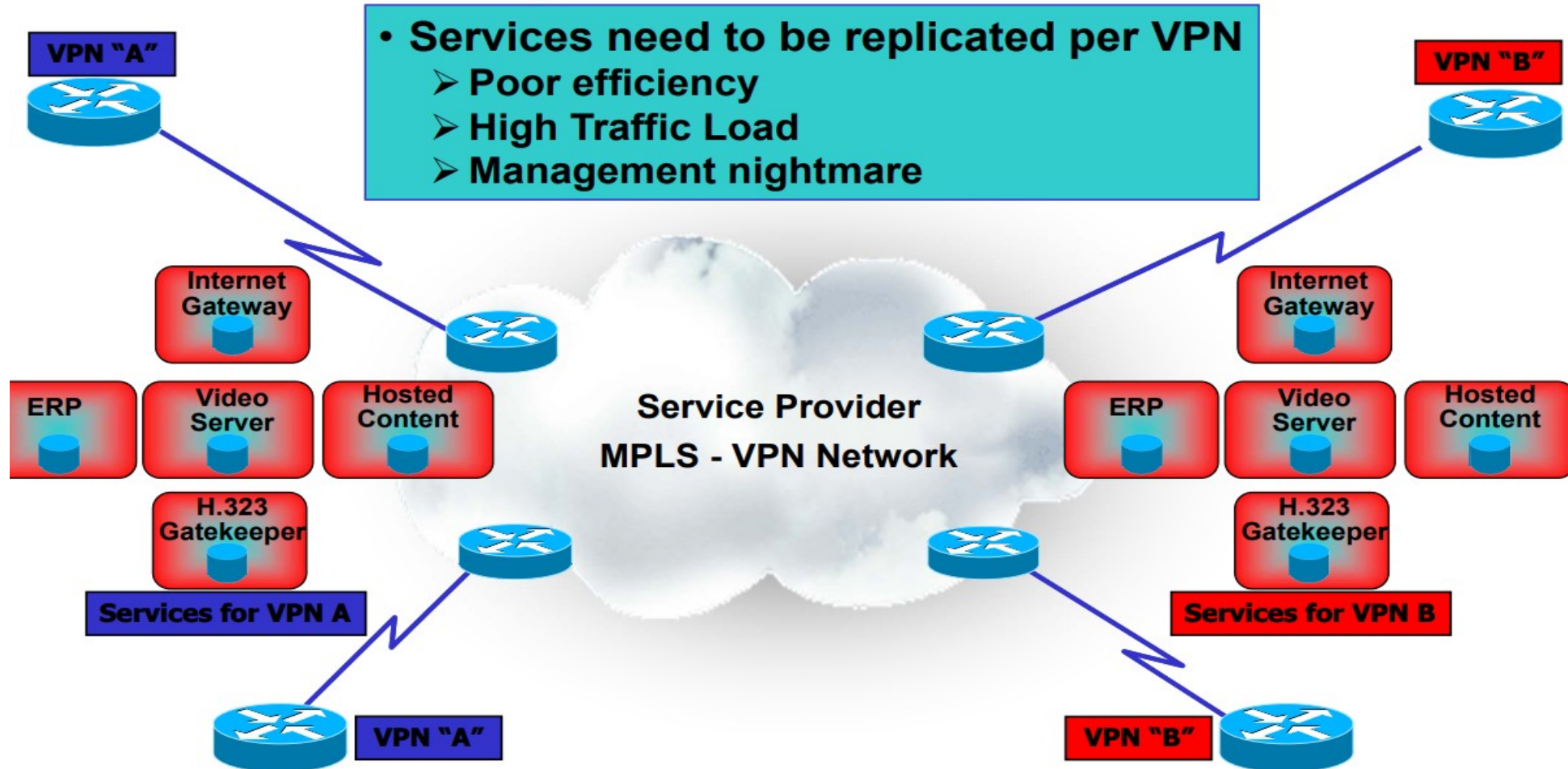
- Un environnement de marché plus délicat
 - Allongement du cycle de vente
 - Budgets clients plus limités
- Évolution de la concurrence
 - Concentration/Intégration des acteurs
 - Concurrence des opérateurs globaux
 - Pression concurrentielle en augmentation
- Préservation de la marge
 - Marge sur le matériel en baisse
- Évolution des métiers
 - Recherche de relais de croissances : VPN IP, VoIP, ToIP, PBXIP, WiFi
 - Risques de désintermédiation
 - En route vers plus de service : Service Management, Infogérance..
 - Maîtrise de la relation client

Contraintes opérateur et entreprise

- **Côté opérateur**
 - Une solution technique permettant d'intégrer des nouvelles technologies d'accès (émergence du DSL) et intégrant les évolutions des cœurs de réseau opérateurs (Gigabit Ethernet...)
 - Topologies d'interconnexion de plus en plus complexes et changeantes au sein d'entreprises plus ouvertes
- **Entreprises** à la recherche de solutions leur offrant :
 - Coûts
 - Sécurité
 - Bonnes performances pour leurs applications métiers
 - Evolutivité de la solution

Avant MPLS/VPN:

Managed Shared Services



Conclusion

- ❑ MPLS est adapté aux besoins du moment
 - ❑ VPN,
 - ❑ QOS,
 - ❑ TE
- ❑ MPLS se base sur l'existant (protocoles) et permet les évolutions futures possibles (IPV6)
- ❑ MPLS fait partie d'un mouvement d'ensemble vers les NGN

Travaux pratiques

MATERIELS DE TRAVAUX PRATIQUE :

- Simulateur : paquet tracer ou GNS 3
- Matériels Cisco

TP

- TP1 : Configuration de OSPF Multi zones
- TP2 : Configuration de base de IS-IS
- TP3 : Configuration de IBGP et EBGP
- TP4 : Configuration de base de IP/MPLS
- TP5 : Configuration de VPN MPLS
- TP6 : Configuration de MPLS TE